



**Servicio Nacional  
del Consumidor**



# **Fraude en Medios de Pago en Chile: Determinantes de la Autoprotección del Consumidor**

**SERNAC**

**Subdirección de Consumo Financiero**

Coordinación de Economía del Comportamiento

**Universidad de Chile**

**Facultad de Economía y Negocios**

Departamento de Administración

Santiago, julio de 2025



### **Equipos responsables de la publicación:**

#### **Equipo SERNAC:**

Miguel Pavéz, Subdirector de Consumo Financiero (S);  
Marcela Palominos, Coordinadora Economía del Comportamiento;  
Guillermo Acuña, Analista Senior de la Coordinación de Economía del Comportamiento.

#### **Equipo Facultad de Economía y Negocios de la Universidad de Chile:**

Cristóbal Barra, profesor asistente de la Facultad de Economía y Negocios de la Universidad de Chile, Ph.D. in Marketing, University of South Carolina.  
Ignacio Vargas, ayudante de Investigación Doctoral de la Facultad de Economía y Negocios de la Universidad de Chile, Ph.D.(c) in Marketing, University of North Texas.

SERNAC agradece la valiosa contribución en la realización de este estudio a Andrés Pavón (Ex Subdirector de Consumo Financiero, MSc Regulation, London School of Economics, MSc Public Policy, University College London); Carlos Noton (profesor asistente de la Escuela de Administración de la Universidad Católica de Chile, PhD in Economics, University of California, Berkeley).



## Contenido

|                                           |    |
|-------------------------------------------|----|
| <b>1. Introducción</b>                    | 4  |
| <b>2. Marco Teórico</b>                   | 6  |
| 2.1 Teoría de la motivación de protección | 6  |
| 2.2 Teoría del comportamiento planificado | 8  |
| <b>3. Diseño Metodológico</b>             | 10 |
| 3.1 Modelo SEM                            | 10 |
| 3.2 Variables                             | 11 |
| 3.3 Procedimientos de muestreo            | 15 |
| <b>4. Resultados</b>                      | 17 |
| 4.1 Caracterización de la Muestra         | 17 |
| 4.2 Modelo de Medición                    | 22 |
| 4.3 Modelo Estructural                    | 25 |
| 4.4 Análisis de Robustez                  | 27 |
| 4.5 Análisis de Heterogeneidad            | 29 |
| <b>5. Conclusiones</b>                    | 32 |
| <b>6. Bibliografía</b>                    | 34 |
| <b>ANEXO 1: Modelo SEM</b>                | 39 |
| <b>ANEXO 2: Variables Latentes</b>        | 51 |



## **1. Introducción**

La reciente promulgación de la Ley N° 21.673, destinada a combatir el sobreendeudamiento y publicada en el Diario Oficial el 30 de mayo de 2024, reformuló el marco regulatorio relacionado a la limitación de la responsabilidad de los usuarios de medios de pago con ocasión del fraude, contenido en la Ley N° 20.009. La nueva normativa impone nuevas obligaciones, tanto a los usuarios de medios de pagos, como a las entidades reguladas. Por una parte, la reforma legal dispuso que “los usuarios deberán informarse y adoptar todas las medidas necesarias para prevenir el uso indebido, el fraude u otros riesgos afines a la utilización de los medios de pago a que se refiere esta ley y los mecanismos de autenticación asociados”. Paralelamente, las entidades reguladas tienen el deber de proveer información periódica, clara, accesible y actualizada sobre las medidas de seguridad y las instrucciones para un uso seguro, fomentando prácticas responsables en la gestión de los medios de pago (art. 4 bis).

La reforma normativa destaca la importancia del componente humano dentro de la seguridad informática. Si bien la tecnología es indispensable, la literatura especializada coincide en que su eficacia es limitada sin la promoción activa de conductas de autoprotección entre los consumidores (Furnell et al., 2006; Liang & Yajiong, 2010; Ng Boon-Yuen et al., 2009; Rhee et al., 2009). En este contexto, las ciencias del comportamiento emergen como herramientas fundamentales para el diseño de políticas de prevención de fraudes más robustas y efectivas (Anderson et al., 2010; Acquisti et al., 2015). La adopción de medidas preventivas por parte de los consumidores representa una estrategia clave para mitigar los riesgos de fraude (Jansen & Van Schaik, 2018b).

La seguridad de la banca en línea depende significativamente del comportamiento del usuario, ya que los bancos tienen limitaciones para controlar las acciones de sus clientes y los dispositivos que estos utilizan. De hecho, estudios como el de Jansen y Leukfeldt (2015) han demostrado que el comportamiento del usuario es un factor crítico en la victimización por fraude bancario en línea. Por lo tanto, es imperativo que los clientes estén alerta y preparados para enfrentar las amenazas dirigidas a las plataformas de banca en línea. Esto implica desarrollar una conciencia aguda de los riesgos, implementar medidas preventivas y saber cómo reaccionar ante incidentes de seguridad. En este contexto, una estrategia fundamental es fomentar la adopción de prácticas de precaución entre los usuarios.

Por consiguiente, esta investigación se enfoca en identificar los factores que motivan a los consumidores a adoptar comportamientos preventivos. Un entendimiento profundo de estas motivaciones permitirá establecer los pilares para el desarrollo de incentivos informativos eficaces, destinados a educar a los consumidores sobre los riesgos de fraude y a fomentar prácticas seguras.

El presente estudio evaluó empíricamente las principales teorías conductuales predictivas de la motivación a la prevención del fraude. Para ello, se analizan los resultados de una encuesta aplicada a más de 2.000 consumidores, utilizando un modelo de ecuaciones estructurales (SEM, por sus siglas en inglés), que integra dos métodos estadísticos: el Modelo de Medición y el Modelo Estructural. Este estudio evaluó empíricamente la aplicabilidad de la Teoría de la Motivación de Protección (TMP; Rogers,

1975; Maddux y Rogers, 1983) y la Teoría del Comportamiento Planificado (TCP; Ajzen, 1991) al contexto específico de la prevención del fraude financiero en línea. La TMP, reconocida por su capacidad predictiva en relación con conductas preventivas (Floyd et al., 2000), se centra en el análisis de la toma de decisiones ante situaciones de riesgo. En contraparte, la TCP proporciona un marco conceptual para comprender la influencia de factores internos y externos sobre las intenciones conductuales, lo que la convierte en una herramienta valiosa para el diseño de intervenciones orientadas al cambio de comportamiento (Yousafzai et al., 2010; Brown et al., 2016). Los hallazgos del estudio demuestran que la integración de ambas teorías ofrece una comprensión significativamente más completa de los determinantes de la intención de protección contra el fraude en línea, en comparación con la aplicación individual de cada modelo.

Los resultados del estudio mostraron que el factor más importante para motivar la protección contra el fraude es la denominada "autoeficacia", es decir, la creencia de una persona sobre su capacidad para realizar la conducta que se le recomienda en materia de seguridad contra el fraude online. Por otro lado, aunque menos robusta, se encuentra que la percepción de la "gravedad" del fraude también juega un papel esencial. Los consumidores que consideran que las consecuencias de ser víctimas de fraude son graves están más motivados para adoptar medidas preventivas. Esto sugiere que la comunicación eficaz sobre los riesgos y consecuencias del fraude puede ser una herramienta poderosa para fomentar comportamientos protectores. Adicionalmente, los "costos" percibidos de la respuesta, tanto en términos de tiempo, dinero como esfuerzo, tienen un impacto negativo en la motivación de protección. La reducción de estos costos percibidos, mediante la simplificación de las medidas de seguridad y la automatización de procesos de protección, puede aumentar la disposición de los consumidores a seguir las recomendaciones de seguridad.

Este estudio, basado en evidencia empírica sólida, ofrece información valiosa para la política pública. Para motivar a los consumidores a protegerse del fraude informático, se sugiere que las intervenciones preventivas se enfoquen en fortalecer su sentido de autoeficacia. Esto implica empoderar a las personas para que se sientan capaces de salvaguardar sus sistemas de información. Los resultados de esta investigación pueden guiar a instituciones financieras y organismos públicos en el diseño de campañas educativas, de formación y concientización sobre ciberseguridad. El objetivo es promover el uso seguro de los medios de pago en línea, preparando mejor a los individuos para enfrentar y mitigar las amenazas digitales.

La estructura de este estudio se presenta de la siguiente manera: Tras la introducción, la Sección 2 (Marco Teórico) describe las dos teorías conductuales centrales analizadas: la Teoría de la Motivación de Protección y la Teoría del Comportamiento Planificado. La Sección 3 detalla el diseño metodológico, comenzando con la descripción del método de análisis (Modelado de Ecuaciones Estructurales - SEM) y sus etapas secuenciales, específicamente las etapas del modelo de medición y el modelo estructural. A continuación, se presentan las variables incluidas en el análisis (dependientes, independientes y de control), así como los procedimientos de muestreo utilizados. En la Sección 4 se exponen los resultados principales, iniciando con la caracterización de la muestra, seguida de los resultados del modelo de medición y del modelo estructural. También se incluyen los resultados del análisis de robustez, considerando la inclusión de controles mediante Mínimos Cuadrados Ordinarios (MCO) y regresión de cuantiles. Además, se presentan los resultados del análisis de heterogeneidad, donde se evalúan

las diferencias entre subgrupos basados en género, edad, ingreso, educación, experiencia, conocimiento y experiencias previas de fraude. Finalmente, la Sección 5 ofrece las conclusiones principales del estudio.

## **2. Marco Teórico**

Para estudiar los factores que motivan a los consumidores a protegerse contra el fraude informático se consideraron dos teorías conductuales: La teoría de la Motivación de Protección (TMP) (Rogers, 1975; Maddux & Rogers, 1983); y la Teoría del Comportamiento Planificado (TCP) (Yousafzai, et al., 2010). Ambas teorías se describen a continuación.

### **2.1 Teoría de la motivación de protección**

La Teoría de la Motivación de Protección (PMT, por sus siglas en inglés) es un modelo teórico desarrollado por Rogers (1975) que se utiliza para explicar cómo las personas toman decisiones relacionadas con la protección en situaciones de riesgo (Rogers, 1975; Maddux & Rogers, 1983). Aunque en un principio fue formulada para comprender las conductas asociadas a la salud, su aplicación se ha extendido a ámbitos como la seguridad pública y la ciberseguridad, entre otros (Maddux & Rogers, 1983; Milne et al., 2000; Floyd et al., 2000).

La PMT sostiene que la respuesta de una persona ante una amenaza depende de dos procesos principales (Rogers, 1975; Maddux & Rogers, 1983):

- (1) **Evaluación de la amenaza:** El individuo evalúa la probabilidad y el impacto de una amenaza. Esto incluye la percepción de la vulnerabilidad personal ante esa amenaza y la percepción de la gravedad de una amenaza potencial. Si una persona considera que la amenaza es grave y se siente vulnerable a ella, es más probable que se sienta motivada para tomar medidas.
- (2) **Evaluación del afrontamiento:** Una vez que la persona reconoce la amenaza y la evalúa, inicia un proceso en el que analiza las estrategias disponibles para mitigarla. Este proceso incluye la percepción de la eficacia de la respuesta recomendada, la autoeficacia (es decir, la confianza en la propia capacidad para ejecutar la acción propuesta) y la evaluación de los costos asociados a la respuesta. La interacción de estos factores determina si una persona percibe la acción de protección como útil y si se siente capaz de llevarla a cabo.

Finalmente, tras evaluar la amenaza y las estrategias de afrontamiento, los individuos adoptan un comportamiento adaptativo o no adaptativo. Los comportamientos adaptativos corresponden a acciones destinadas a la protección frente a la amenaza, mientras que las respuestas no adaptativas incluyen aquellas situaciones en las que el individuo decide no adoptar medidas de protección.



Tanto en el modelo teórico como en el empírico, la motivación de protección es la variable dependiente, entendida como la intención de iniciar, mantener o evitar ciertos comportamientos para prevenir fraudes financieros (Floyd et al., 2000). Específicamente, se refiere a la intención de seguir las recomendaciones de seguridad contra el fraude informático (es decir, la intención de protegerse), las cuales suelen ser proporcionadas por las instituciones financieras para que los consumidores se resguarden ante estas amenazas.

Por otro lado, de acuerdo con diversos estudios (Liang & Yajiong, 2010; Jansen & Van Schaik, 2018b), las variables independientes o predictoras del modelo, adaptadas al contexto de la banca en línea, se derivan de dos procesos clave: la evaluación de la amenaza de fraude y la evaluación del afrontamiento del fraude, realizados por los consumidores. A continuación, se describen estas variables.

### **Evaluación de la amenaza**

Este proceso consiste en la valoración del nivel de riesgo que un individuo está dispuesto a asumir (Workman et al., 2008) y se compone de dos factores:

**Vulnerabilidad percibida:** Representa la probabilidad de que ocurra un evento adverso si no se toman medidas de protección adecuadas (Rogers, 1975). En otras palabras, en el contexto de este estudio, es la percepción de la probabilidad de ser víctima de un fraude en la banca en línea (Crossler, 2010).

**Gravedad percibida:** Se refiere a la magnitud de las consecuencias que podría generar dicho evento de seguridad (Rogers, 1975), es decir, la evaluación subjetiva del usuario sobre el impacto de ser víctima de un fraude informático (Crossler, 2010).

La Teoría de la Motivación de Protección (TMP) postula que tanto una mayor percepción de vulnerabilidad como una mayor gravedad percibida incrementan la probabilidad de que un consumidor adopte medidas de protección contra el fraude informático.

### **Evaluación del afrontamiento**

Este proceso incluye la valoración de estrategias y recursos disponibles para mitigar la amenaza. Sus principales componentes son:

**Eficacia de la respuesta:** Es la creencia de que una determinada acción producirá el resultado esperado (Rogers, 1975). En este contexto, se refiere a la percepción de los usuarios sobre la efectividad de una medida de seguridad para reducir el riesgo de fraude (Milne et al., 2000). Si una persona confía en que una estrategia de protección es eficaz, será más propensa a adoptarla.

**Autoeficacia:** Es la confianza del individuo en su capacidad para ejecutar una acción específica (Maddux & Rogers, 1983). En el ámbito de la ciberseguridad, se traduce en la creencia de un usuario en su habilidad para proteger sus sistemas de información (Rhee



et al., 2009). Una mayor autoeficacia está asociada a una mayor adopción de medidas de protección.

**Costos de la respuesta:** Se refieren a los recursos percibidos como necesarios para implementar una medida de seguridad (Milne et al., 2000). Estos costos pueden ser de diversa naturaleza, incluyendo tiempo, dinero o esfuerzo (Menard et al., 2017). Las medidas de protección solo serán adoptadas si los costos percibidos son menores que las posibles pérdidas derivadas de la amenaza (Crossler, 2010).

Es importante destacar que estos componentes corresponden a procesos cognitivos y, por lo tanto, son constructos psicológicos. Es decir, son variables teóricas y abstractas utilizadas para explicar el comportamiento de los individuos. Aunque no son directamente observables, pueden medirse de manera indirecta a través de indicadores específicos (Slaney & Racine, 2013). Estas mediciones suelen realizarse mediante ítems, que pueden presentarse en forma de preguntas, afirmaciones o tareas diseñadas para evaluar distintos aspectos de un constructo psicológico.

**Tabla 1: Procesos cognitivos de la TMP y sus componentes como explicación del comportamiento precautorio**

| <b>Evaluación de la amenaza</b>                                                                                                  | <b>Evaluación del afrontamiento</b>                                                                                                                                  |
|----------------------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <ul style="list-style-type: none"> <li>▪ <i>Vulnerabilidad percibida (+)</i></li> <li>▪ <i>Gravedad percibida (+)</i></li> </ul> | <ul style="list-style-type: none"> <li>▪ <i>Eficacia de la respuesta (+)</i></li> <li>▪ <i>Autoeficacia (+)</i></li> <li>▪ <i>Costos de respuesta (-)</i></li> </ul> |

Entre paréntesis se muestra el signo de la relación entre cada componente y el comportamiento precautorio.

## **2.2 Teoría del comportamiento planificado**

Otro marco teórico relevante para este estudio es la Teoría del Comportamiento Planificado (TCP), propuesta por Ajzen (1991, 2002). Esta teoría se ha utilizado ampliamente para explicar el comportamiento humano en diversos ámbitos. Por ejemplo, ha sido aplicada para predecir prácticas deshonestas en estudiantes universitarios (Mayhew et al., 2009), comprender la intención de utilizar la banca en línea (Shih & Fang, 2004; Yousafzai et al., 2010), estudiar la intención de los contadores profesionales de denunciar fraudes contables (Brown et al., 2016) y analizar comportamientos de protección en línea (Burns & Roberts, 2013), entre otros.

La TCP postula que el comportamiento humano está determinado por tres tipos de creencias fundamentales (Ajzen, 1991, 2002):

**Creencias conductuales (Actitud):** Se refieren a las creencias sobre las posibles consecuencias o atributos de un comportamiento, lo que determina la actitud hacia el comportamiento. Es decir, la percepción positiva o negativa que una persona tiene sobre realizar una acción específica. Esta actitud se basa en la evaluación de los resultados esperados al llevar a cabo el comportamiento.



**Creencias normativas (Norma Subjetiva):** Son las creencias sobre las expectativas de otras personas significativas en el entorno del individuo, lo que da lugar a las normas subjetivas. Estas reflejan la percepción de la presión social para realizar o no un comportamiento y están influenciadas por la opinión de familiares, amigos, colegas u otros grupos de referencia.

**Creencias de control (Control conductual percibido):** Se refieren a la percepción sobre la presencia de factores que pueden facilitar o dificultar la realización del comportamiento, lo que determina el control conductual percibido. Este aspecto implica la evaluación de la facilidad o dificultad para llevar a cabo una acción, basada en experiencias previas y posibles obstáculos. Es un factor clave, ya que puede reforzar o debilitar la intención de actuar, dependiendo de cuán capacitado se sienta el individuo para ejecutar el comportamiento.

En conjunto, estos tres elementos—actitud hacia el comportamiento, norma subjetiva y control conductual percibido—contribuyen a la formación de la intención de comportamiento. La TCP asume que la intención es el principal antecedente del comportamiento y que, cuando las condiciones lo permiten, las personas actúan en función de sus intenciones (Ajzen, 2002).

Por ejemplo, en el contexto de este estudio: i) Una actitud favorable hacia la seguridad financiera implica considerar que seguir las recomendaciones es correcto, útil o beneficioso; ii) Una norma subjetiva favorable significa percibir que las personas importantes para el individuo (como familiares, amigos o colegas) esperan que siga las recomendaciones de seguridad; iii) Un control conductual percibido alto se traduce en la creencia de que seguir las recomendaciones de seguridad es fácil y factible.

Cuando una persona tiene una actitud positiva hacia un comportamiento, percibe que los demás lo consideran apropiado y cree que es fácil de realizar, la probabilidad de que lo adopte aumenta significativamente.

Finalmente, es importante señalar, que el constructo control conductual percibido ha sido relacionado con el concepto de autoeficacia, desarrollado por Bandura (1977) y retomado por Ajzen (2002). La autoeficacia se define como la creencia de una persona en su capacidad para realizar un comportamiento específico. Aunque ambos conceptos son parecidos, no son idénticos.

Por una parte, la autoeficacia se enfoca en la percepción del individuo sobre su capacidad personal para realizar una acción. En tanto, el control conductual percibido se centra en la percepción de la facilidad o dificultad del comportamiento en sí, considerando factores internos y externos que pueden influir en su ejecución.

Algunos estudios han tratado estos conceptos como equivalentes, especialmente en investigaciones que combinan la Teoría de la Motivación de Protección (TMP) y la Teoría del Comportamiento Planificado (TCP). Por ejemplo, Ifinedo (2012) analizó comportamientos de protección y trató la autoeficacia como un componente común a ambas teorías.



En este estudio, no se asumió a priori que la autoeficacia y el control conductual percibido fueran el mismo concepto. En su lugar, se trataron como constructos diferenciados y, posteriormente, en el análisis empírico se evaluó la conveniencia de unificarlos en un solo factor.

### **3. Diseño Metodológico**

#### **3.1 Modelo SEM**

En este estudio se empleó el Modelado de Ecuaciones Estructurales (SEM, por sus siglas en inglés), un conjunto de técnicas estadísticas multivariadas diseñadas para modelar relaciones complejas entre variables observadas y latentes (Hoyle, 2012). El SEM permite analizar simultáneamente múltiples relaciones entre variables, integrando tanto variables observadas (también llamadas manifiestas) como variables latentes, es decir, variables que no son directamente medibles, pero que se infieren a partir de múltiples indicadores observables (Pett et al., 2003).

El SEM modela estas relaciones mediante ecuaciones estructurales, que representan conexiones hipotéticas, frecuentemente de carácter causal o secuencial, entre las variables. Su uso se orienta a la evaluación de modelos teóricos: el investigador plantea un modelo basado en hipótesis derivadas de la teoría, y SEM permite verificar empíricamente si los datos respaldan dicho modelo (Ramlall, 2016). Por esta razón, el SEM se considera una técnica confirmatoria, más que exploratoria; es decir, su propósito principal es confirmar o refutar hipótesis preexistentes, en lugar de identificar relaciones emergentes sin guía teórica.

El enfoque SEM combina dos componentes fundamentales: a) El modelo de medición, que especifica qué variables observadas (ítems) miden cada factor latente y los patrones de covarianza entre estos ítems; b) El modelo estructural, que especifica las relaciones causales o correlacionales entre las variables latentes, diferenciando entre variables independientes y dependientes (Ramlall, 2016).

En el Anexo 1, se describen brevemente los distintos aspectos del Modelo de Medición y del Modelo Estructural. En el caso específico de este estudio, el Modelo de Medición incluyó el Análisis Factorial Exploratorio, Análisis de consistencia interna, Análisis Factorial Confirmatorio y Análisis de validez convergente y discriminante.

## 3.2 Variables

### Variable Dependiente

La variable dependiente en este estudio es la **motivación de protección**, definida como la intención de seguir las recomendaciones de seguridad contra el fraude informático. Estas recomendaciones corresponden a las medidas de protección que las instituciones financieras suelen proporcionar a sus clientes para reducir los riesgos asociados a la banca en línea.

La elección de medir la **intención de seguir las recomendaciones de seguridad** se justifica porque, según diversos estudios, la intención es el **predictor más inmediato del comportamiento real** (Fishbein & Ajzen, 2010). Además, se ha demostrado que es posible modificar el comportamiento de las personas influyendo en sus intenciones (Anderson et al., 2010; O’Keefe, 2016).

Para determinar qué recomendaciones incluir en el estudio, se llevó a cabo un estudio cualitativo basado en la recopilación y análisis de las directrices de seguridad emitidas por instituciones financieras en Chile durante febrero de 2024. Este análisis permitió identificar las recomendaciones más frecuentes y formular versiones representativas de cada una.

Como resultado, se obtuvo la siguiente lista de recomendaciones (ver Tabla 2).

**Tabla 2: Lista de recomendaciones representativas**

1. Evite abrir enlaces y descargar archivos adjuntos. Su institución financiera nunca le pedirá este tipo de acciones.
2. Entre siempre a su institución financiera ingresando la dirección en la barra del navegador. Nunca ingrese mediante enlaces en correos electrónicos o buscadores como Google.
3. Al conectarse a Internet, hágalo desde dispositivos propios, y nunca mediante una red pública y abierta.
4. Mantenga en privado sus claves. No las comparta, ni siquiera con familiares ni amigos.
5. Evite guardar y enviar fotos de su tarjeta.
6. Cambie sus claves frecuentemente.
7. Memorice sus claves. No permita que se guarden en su computador o celular.
8. Use claves complejas, que no sean fáciles de deducir como fechas de nacimiento o números de teléfono.

Entonces, Para medir la **intención de protegerse (variable dependiente)**, se preguntó a los participantes, para cada una de las recomendaciones de seguridad presentadas: “¿Qué tan probable es que siga las siguientes recomendaciones?”.

Las respuestas se registraron en una escala de Likert, codificadas numéricamente de 1 (muy improbable) a 5 (muy probable). De esta manera, la variable dependiente refleja la intención de seguir un conjunto de recomendaciones de seguridad, lo cual resulta

adecuado, dado que el comportamiento precautorio frente a amenazas requiere la ejecución de múltiples acciones (Crossler & Bélanger, 2014).

## **Variables independientes**

Este estudio exploró la capacidad predictiva de la Teoría de la Motivación de Protección (TMP) y la Teoría del Comportamiento Planificado (TCP), adaptadas al contexto específico de nuestra investigación. Se consideraron inicialmente ocho constructos psicológicos: vulnerabilidad percibida, gravedad percibida, autoeficacia, eficacia de la respuesta, costos de la respuesta, actitud hacia el comportamiento, normas subjetivas y control conductual percibido.

Para la medición de estos constructos, se empleó una escala Likert de 5 puntos, donde 1 representaba "totalmente en desacuerdo" y 5 "totalmente de acuerdo". Cada afirmación fue evaluada utilizando esta escala. En la fase inicial del estudio, se utilizó un conjunto de seis ítems por constructo, adaptados de escalas previamente validadas en investigaciones anteriores. A continuación, se detallan los constructos considerados inicialmente en el modelo de medición. A la vez, en el anexo 2, se presentan los diferentes ítems utilizados inicialmente para la definición de cada constructo.

### **Teoría de Motivación de Protección**

- **Evaluación de amenazas a la seguridad:** La evaluación de un individuo sobre el nivel de peligro que representa un evento de seguridad (Crossler, 2010).

**Vulnerabilidad de seguridad percibida:** La vulnerabilidad percibida es la evaluación que hace un usuario sobre la probabilidad de que ocurra un evento de seguridad amenazante (Crossler, 2010). En este caso, se refiere a la medida en que un cliente cree que será víctima de fraude en la banca en línea.

**Gravedad percibida:** es la evaluación subjetiva de un usuario sobre la magnitud de las consecuencias negativas derivadas de un evento de seguridad amenazante. En el contexto de la banca en línea, este estudio la aplica a la percepción de un cliente sobre la seriedad de las consecuencias del fraude (Crossler, 2010).

- **Evaluación del afrontamiento de seguridad:** La evaluación de una persona sobre su capacidad para realizar una acción específica y su confianza en que esa acción evitará pérdidas o daños ante una amenaza de seguridad, considerando que el costo percibido no sea excesivo (Crossler, 2010).

**Autoeficacia:** La confianza de un individuo en su capacidad para realizar el comportamiento recomendado para prevenir o mitigar el evento de seguridad amenazante (Crossler, 2010).

**Eficacia de la respuesta:** Se define como la percepción de una persona sobre la efectividad de una acción para reducir o eliminar una amenaza. En este estudio, se

refiere a la confianza en que un comportamiento recomendado prevendrá o mitigará un evento de seguridad amenazante (Crossler, 2010).

**Costos de la respuesta:** Se definen como la percepción del usuario sobre el esfuerzo cognitivo, tiempo, dinero u otros recursos necesarios para implementar una acción de protección (Milne et al., 2000; Menard et al., 2017). Según Crossler (2010), las medidas de precaución se adoptan cuando sus costos percibidos son menores que las posibles pérdidas por una amenaza. Es decir, este factor enfatiza los costes de oportunidad percibidos en términos de dinero, tiempo y esfuerzo invertidos en adoptar el comportamiento recomendado. Por lo tanto, el costo percibido influye directamente en la disposición del usuario a actuar ante riesgos de seguridad de la información.

### **Teoría del Comportamiento Planificado**

**Normas subjetivas:** Definen la percepción de un individuo sobre las creencias de las personas relevantes para él en relación con un comportamiento específico, en este caso, tomar medidas de seguridad contra el fraude informático (Ifinedo, 2012). Esencialmente, reflejan la creencia del individuo sobre la aprobación o desaprobación de otros significativos hacia su participación en dicha conducta.

**Actitud:** Se refiere al grado en que una persona evalúa favorable o desfavorablemente un comportamiento de interés, en este caso, seguir las recomendaciones de seguridad, considerando los resultados de llevarlas a cabo. En otras palabras, son los sentimientos positivos o negativos del individuo hacia la participación en esta conducta específica (Ifinedo, 2012).

**Control conductual percibido:** Se refiere a la percepción de una persona sobre la facilidad o dificultad de realizar el comportamiento de interés. Según Ifinedo (2012), esta percepción varía según las situaciones y acciones, lo que resulta en que una persona tenga percepciones variables del control conductual según la situación. Mientras que el énfasis de la autoeficacia está en si los individuos sienten que tienen las habilidades y capacidades adecuadas para lograr un objetivo, el control conductual percibido comprende una expresión más interactiva de la relación entre un individuo y su entorno (Workman et al., 2008).

### **Variables de Control**

En las pruebas de robustez se utilizaron las siguientes variables de control:

**Informado:** La variable "Informado" evalúa el nivel de autopercepción de los encuestados sobre su conocimiento respecto a los riesgos de fraude informático. Para ello, se planteó la pregunta: "¿Qué tan bien informado se siente sobre los riesgos del fraude informático?". Las respuestas se midieron en una escala Likert de 5 puntos, donde 1 representa "Nada informado" y 5 "Muy bien informado", con las siguientes opciones intermedias: 2) No muy bien informado; 3) Algo informado; 4) Bastante bien informado.



**Experiencia:** La variable "Experiencia" mide la autopercepción de los encuestados respecto a su nivel de habilidad en el uso del sitio web o aplicaciones de su institución financiera. Para ello, se preguntó: "¿Se considera experimentado(a) en el uso del sitio web o aplicaciones de su institución financiera?". Las respuestas se presentan en una escala de cinco niveles, desde "Muy por debajo del promedio", que indica una baja autoevaluación, hasta "Muy por encima del promedio", que señala una alta autoevaluación.

**Fraude:** La variable "Fraude" registra la victimización de los encuestados por fraude informático en la banca en línea durante el último año. Se preguntó directamente: "¿Ha sido víctima de fraude informático en la banca en línea durante el último año? Esto incluye, por ejemplo, cobros desconocidos en su tarjeta de crédito o cuenta bancaria, transferencias no reconocidas, clonación de tarjetas o hackeo de cuentas". Las respuestas posibles fueron: "No", "Sí, una vez" y "Sí, más de una vez". Esta variable mide la incidencia directa de experiencias de fraude informático en los participantes del estudio.

**Género:** La variable "Genero" identifica el género con el que se identifica el encuestado. Las opciones son "Femenino", "Masculino", "Otro" y "Prefiero no decirlo". Esta variable permite clasificar a los participantes según su identidad de género, ofreciendo la opción de no revelar esta información si así lo desean.

**Edad:** La variable "Edad" clasifica a los encuestados en grupos de edad. Las categorías son "Entre 18 y 29 años", "Entre 30 y 44 años", "Entre 45 y 59 años" y "60 años o más".

**Educación:** La variable "Educación" clasifica el nivel educativo máximo alcanzado por los participantes. Las categorías se organizan en orden ascendente de formación académica, desde "No estudió o educación básica incompleta" hasta "Educación de postgrado (pos título, master, magíster, doctor)". Las opciones de respuesta fueron: 1) No estudió o educación básica incompleta; 2) Educación básica completa; 3) Educación media completa; 4) Educación superior técnica completa (CFT o IP); 5) Educación universitaria completa; y 6) Educación de postgrado (pos título, master, magíster, doctor).

**Ingreso:** La variable "Ingreso" clasifica a los participantes según su rango de ingreso mensual total, considerando todas las fuentes de ingresos (salarios, alquileres, jubilaciones, pensiones, etc.). Los rangos de ingreso evaluados fueron: 1) Menos de \$250.000; 2) Entre \$250.000 y \$500.000; 3) Entre \$500.001 y \$900.000; 4) Entre \$900.001 y \$1.500.000; 5) Entre \$1.500.001 y \$2.000.000; y 6) Más de \$2.000.000. Este desglose permite analizar las respuestas en relación con el nivel socioeconómico de los encuestados.



### 3.3 Procedimientos de muestreo

#### **Criterios de elegibilidad**

Los criterios de elegibilidad para este estudio fueron: ser cliente de una institución financiera y utilizar personalmente las plataformas de la banca en línea. Además, con el fin de garantizar la calidad de los datos recopilados, se excluyeron las encuestas incompletas o aquellas cuya duración fue inferior al 50% de la mediana del tiempo de respuesta.

#### **Selección de la muestra**

Los participantes de la encuesta fueron seleccionados aleatoriamente de la base de datos de reclamos del Servicio Nacional del Consumidor (SERNAC). Al seleccionar la muestra, se estratificó por género y edad para garantizar una representación equilibrada. Posteriormente, los participantes fueron contactados a través de un correo electrónico de invitación, el cual incluía un enlace al cuestionario en línea.

#### **Recolección de Datos**

Tras una prueba piloto, se realizaron dos levantamientos adicionales de datos utilizando el mismo cuestionario. El **primer levantamiento** ( $N = 800$ ) tenía como objetivo fue validar el modelo de medición. Se realizó entre el 17 de marzo al 1 de abril de 2024. El **segundo levantamiento** ( $N = 1.223$ ), que se ejecutó entre el 8 y 12 de abril, tuvo como objetivo complementar la muestra para la estimación de una ecuación estructural que identificara los factores explicativos de la motivación de protección de los consumidores. Es decir, para la estimación de la ecuación estructural, se combinaron los datos de ambas encuestas, lo que resultó en una muestra final de  $n = 2.023$  participantes. Solo se consideraron válidas las encuestas completamente respondidas.

#### **Tamaño muestral**

El tamaño muestral se escogió en base a tres criterios: (1) tamaño muestral suficiente para detectar efectos de 0,2 desviaciones estándar a un nivel de significancia de 5% y un poder de 80%, (2) contar con un mínimo de 10 observaciones por constructo, y (3) contar con un mínimo de 5 observaciones por parámetro a estimar (usando como base el modelo con mayor número de parámetros) (Fan, et al., 2016).

#### **Flujo de la Encuesta**

La encuesta constó de varias secciones estructuradas de la siguiente manera:

**Consentimiento informado:** Se informó a los participantes sobre los objetivos del estudio, el carácter voluntario de su participación y la confidencialidad y anonimato de sus respuestas.

**Preguntas de elegibilidad:** Se incluyeron dos preguntas filtro para determinar si el encuestado cumplía con los criterios de inclusión: 1) Ser cliente de una institución financiera; 2) Utilizar personalmente los servicios de banca en línea (es decir, no delegar su uso a terceros).



Servicio Nacional  
del Consumidor



**Medición de variables:** Se formularon preguntas para evaluar la motivación de protección (variable dependiente) y sus variables explicativas, conforme a los constructos de la Teoría de la Motivación de Protección (TMP) y la Teoría del Comportamiento Planificado (TCP).

**Caracterización socioeconómica:** Se recopilieron datos sobre género, edad, nivel educativo y nivel de ingresos de los participantes. Estas variables se usaron como controles en las regresiones y, además, para hacer un análisis de heterogeneidad de los resultados.

A la vez, para mejorar la calidad de las respuestas y minimizar posibles sesgos, se adoptaron las siguientes estrategias:

**Sesgo de orden:** Las preguntas fueron presentadas en orden aleatorio para evitar efectos de secuencia.

**Sesgo de deseabilidad social:** Se garantizó la anonimidad y confidencialidad de las respuestas en las instrucciones. Además, se enfatizó que no existían respuestas correctas ni incorrectas.

**Fatiga del encuestado:** Dado que la validación de los constructos requería un cuestionario extenso, se intercalaron preguntas auxiliares entre los bloques principales. Estas preguntas tenían el propósito de mantener la atención del participante y prevenir respuestas mecánicas. Con este fin, se indagó sobre a) el conocimiento de los encuestados sobre los riesgos del fraude informático; b) su nivel de experiencia en el uso de sitios web de banca en línea; c) la frecuencia de uso de dichas plataformas; d) si habían sido víctimas de fraude (una o más veces); e) si conocían a alguna víctima de fraude (una o más veces). Las dos primeras preguntas se utilizaron como controles en las pruebas de robustez (sección x).





## 4. Resultados

### 4.1 Caracterización de la Muestra

Se analizaron los datos de 2.023 consumidores que cumplieron con los criterios de elegibilidad del estudio: ser clientes de instituciones financieras y utilizar personalmente la plataforma de banca en línea. Además, para garantizar la calidad de las respuestas, solo se consideraron encuestas completamente respondidas y cuya duración fuera igual o superior al 50% de la duración mediana.

El análisis demográfico de los participantes reveló una distribución de género donde 945 se identificaron como mujeres (47%) y 1.051 como hombres (52%), con un 1% (27 individuos) que prefirió no especificar o se identificó con otra opción (**Gráfico 1**). En cuanto a la edad, se observó que el grupo de 18 a 29 años fue el menos participativo, representando solo el 17% del total. Los grupos de mayor edad mostraron una mayor participación, con un 24% entre 30-44 años, un 30% entre 45-59 años, y un 29% con 60 años o más (**Gráfico 2**).

El nivel educacional de los participantes se distribuyó de la siguiente manera: un 20% poseía educación básica o media completa, un 24% tenía educación superior técnica, un 37% contaba con educación universitaria, y un 20% había completado estudios de postgrado (**Gráfico 3**). En cuanto al ingreso de los encuestados, se observó que un 46% percibió ingresos de hasta \$900.000 pesos chilenos, mientras que el 54% restante obtuvo ingresos superiores a dicha cifra (**Gráfico 4**).

Gráfico 1: Participación en el estudio por Género (%)

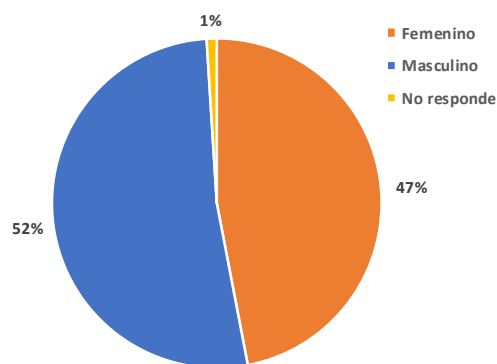
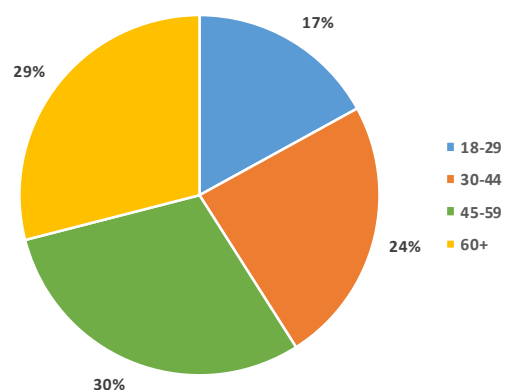
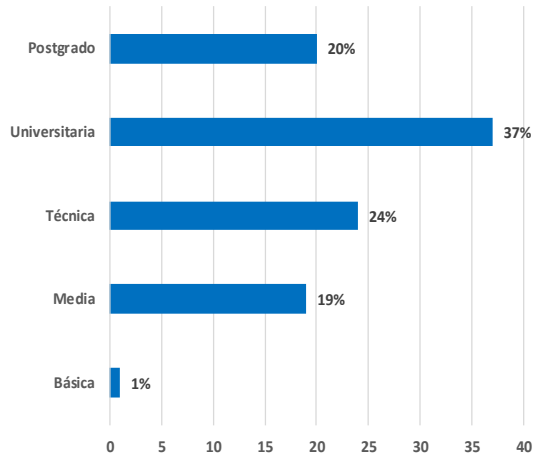


Gráfico 2: Participación en el estudio por rango etario (%)

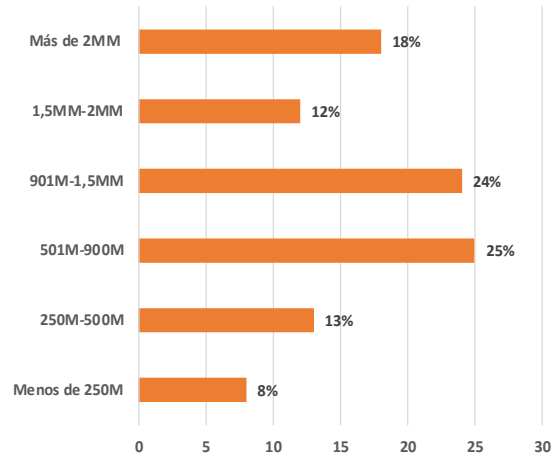


Fuente: Elaboración propia.

**Gráfico 3: Participación en el estudio por nivel educacional (%)**



**Gráfico 4 Participación en el estudio por nivel de ingreso personal (%)**



Fuente: Elaboración propia.

Se consultó a los participantes sobre su percepción de conocimiento acerca de los riesgos del fraude informático, donde el 58% se consideró 'Bastante bien informado' o 'Muy bien informado'. Adicionalmente, se indagó sobre la experiencia de los participantes al utilizar el sitio web o la aplicación móvil de su institución financiera, y el 49% evaluó su experiencia como 'Algo por encima del promedio' o 'Muy por encima del promedio'. Asimismo, el 73% de los participantes declaró utilizar el sitio web o la aplicación móvil de su institución financiera 'Casi todos los días' o 'Todos los días'.

Se consultó a los participantes sobre su experiencia personal con el fraude informático en la banca en línea durante el último año. El 69% declaró no haber sido víctima de fraude. Sin embargo, al preguntarles si conocían a alguien que hubiera sido víctima de fraude informático, el 75% respondió afirmativamente, indicando que conocían a una o más personas que habían experimentado esta situación.

Finalmente, los resultados revelaron que, en relación con la pregunta sobre la victimización por fraude informático en la banca en línea durante el último año, las mujeres experimentaron fraude con mayor frecuencia que los hombres (**Gráfico 5**). De manera similar, el grupo etario de 45 a 59 años mostró la mayor propensión a ser víctima de fraude (**Gráfico 6**). En cuanto al nivel de conocimiento percibido sobre los riesgos del fraude en línea, se observó que los participantes que se autoevaluaron como 'nada informados' o 'no muy bien informados' fueron quienes sufrieron fraude con mayor frecuencia (**Gráfico 7**). Por último, en lo que respecta a la experiencia percibida con la banca en línea, aquellos que se describieron a sí mismos con una experiencia 'algo por debajo' o 'muy por debajo de la media' fueron los más afectados por el fraude informático (**Gráfico 8**).

Gráfico 5: Víctimas de fraude informático en el último año, por género (%)

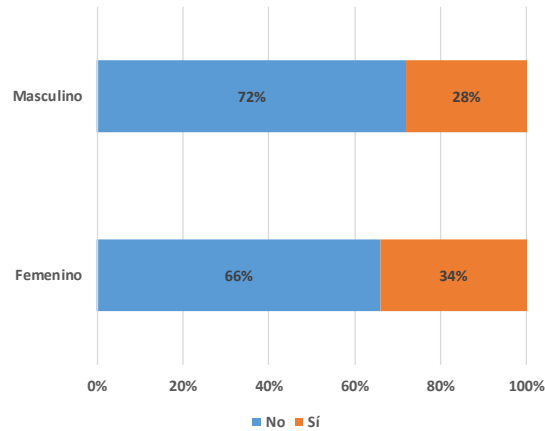


Gráfico 6: Víctimas de fraude informático en el último año, por rango etario (%)

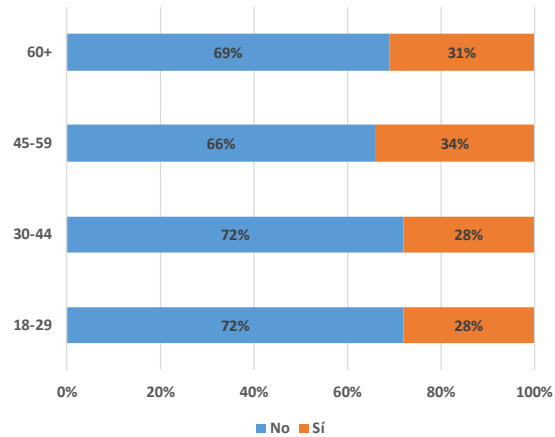


Gráfico 7: Víctimas de fraude informático en el último año, por nivel de conocimiento percibido de los riesgos del ciberfraude (%)

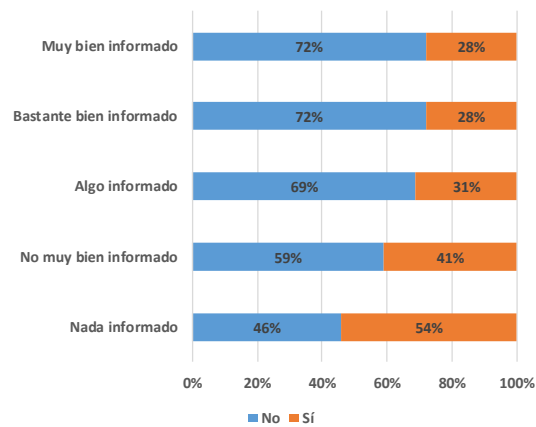
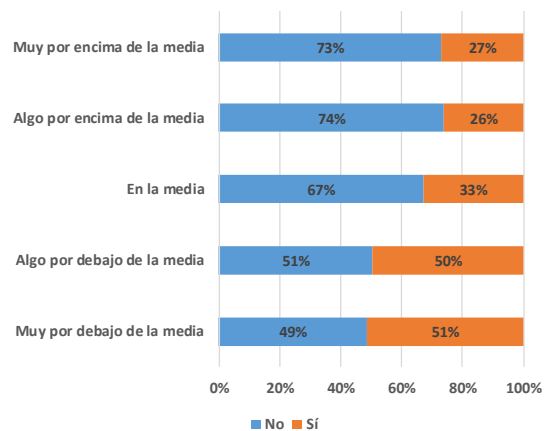


Gráfico 8: Víctimas de fraude informático en el último año, por nivel de experiencia percibida con la banca en línea (%)

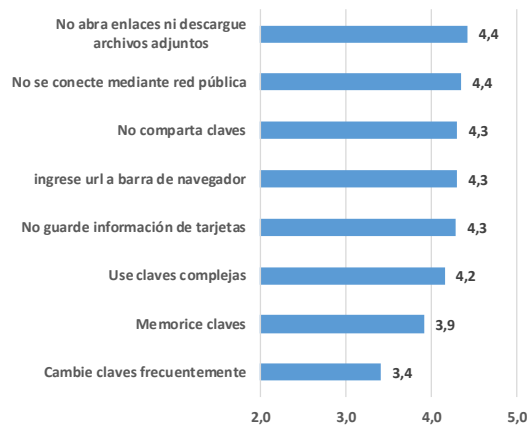


Fuente: Elaboración propia.

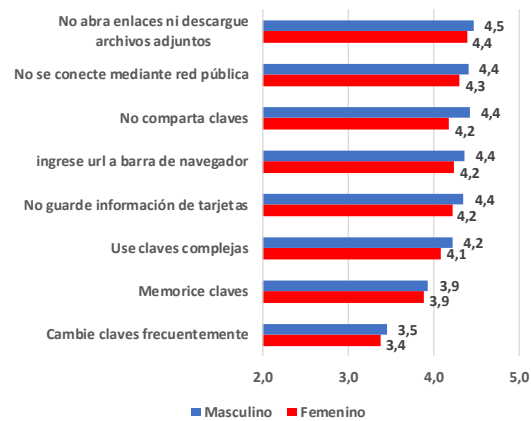
Además, se realizó una regresión para estimar los efectos *ceteris paribus*, con el fin de aislar el impacto de cada variable al mantener constantes las demás. Se utilizó como variable dependiente una variable *dummy* que tomaba el valor de 1 si el participante declaró haber sido víctima de fraude durante el último año, y como variables explicativas, las variables demográficas, de conocimiento y de experiencia. Los coeficientes resultantes se interpretan como el efecto de cada variable sobre la probabilidad de experimentar fraude. Los resultados revelaron dos efectos significativos: la identificación con el género femenino incrementó en un 4% la probabilidad de haber sufrido fraude, mientras que un alto nivel de experiencia con la banca en línea (evaluada como algo o muy por encima del promedio) disminuyó dicha probabilidad en un 8%.

A continuación, se presenta el valor promedio para cada componente de la variable dependiente, la **motivación de protección**, obtenida a partir de la pregunta: '¿Qué tan probable es que siga las siguientes recomendaciones?' (señaladas en la Tabla 2). Los participantes mostraron la mayor probabilidad de seguir la recomendación de 'Evite abrir enlaces y descargar archivos adjuntos...', mientras que la menor probabilidad se observó para 'Cambie sus claves frecuentemente'. Es importante destacar que el valor promedio en todos los casos supera el valor intermedio de 3 (ni probable ni improbable), lo que sugiere que, en promedio, es más probable que improbable que los participantes sigan las recomendaciones de seguridad (**Gráfico 9**).

**Gráfico 9: ¿Qué tan probable es que siga las siguientes indicaciones? (valor promedio)**



**Gráfico 10: ¿Qué tan probable es que siga las siguientes indicaciones? (valor promedio, por género)**



Nota: escala de la pregunta: 1 [muy improbable] a 5 [muy probable]. Se presenta el valor promedio de las respuestas.

Fuente: Elaboración propia.

Un hallazgo relevante fue el valor promedio de cada componente de la variable dependiente, diferenciado por género. En todos los casos, se observó que los hombres reportaron una mayor disposición de seguir las recomendaciones de seguridad (**Gráfico 10**), lo cual podría contribuir a explicar la mayor incidencia de fraudes entre las mujeres en comparación con los hombres.

Para caracterizar las respuestas por **constructo** de la TMP (Teoría de la Motivación de Protección) y la TPB (Teoría del Comportamiento Planificado), se calcularon los promedios de los valores de los ítems correspondientes a cada constructo (mediante un promedio simple) en su escala original. En esta escala, un valor cercano a 5 indica un mayor acuerdo promedio a la afirmación presentada, mientras que un valor cercano a 1 señala desacuerdo. Los resultados muestran que el constructo con mayor acuerdo es '*Gravedad*', lo que sugiere que los participantes consideran que sufrir un fraude informático tendría consecuencias graves. En contraste, el constructo '*Costos*' presenta

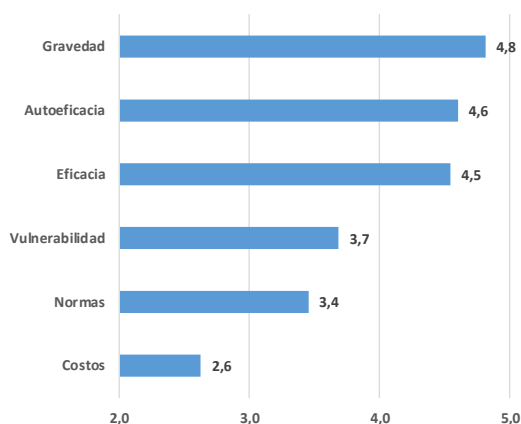


Servicio Nacional  
del Consumidor

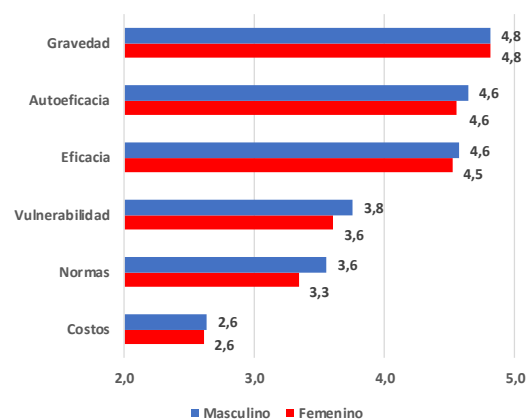
el menor valor promedio, lo que implica que seguir las recomendaciones de seguridad se percibe como poco costoso (**Gráfico 11**).

Al analizar las diferencias por género, se observa que el grado de acuerdo con los ítems de cada constructo tiende a ser mayor entre los hombres que entre las mujeres, lo cual podría explicar la mayor motivación de protección observada en los hombres. La única excepción es '*Gravedad*', donde el grado de acuerdo es ligeramente mayor en las mujeres. Sin embargo, es importante resaltar que, en el caso de '*Costo*', la interpretación del valor es inversa: un menor costo percibido incentiva el comportamiento de protección (**Gráfico 12**).

**Gráfico 11: Valor promedio de cada constructo**



**Gráfico 12: Valor promedio de cada constructo (por género)**



Nota: escala de la pregunta: 1 [muy en desacuerdo] a 5 [muy de acuerdo]. Se presenta el valor promedio de las respuestas.

Fuente: Elaboración propia.



## 4.2 Modelo de Medición

El modelo de medición fue evaluado exclusivamente con los datos del primer levantamiento de la encuesta (N=800). Este enfoque permite una evaluación más rigurosa de la validez del modelo, ya que evita el sobreajuste a los datos, lo que a su vez mejora su validez externa.

Para cada una de las etapas, se usaron los siguientes indicadores y criterios (**Tabla 4**):

**Tabla 4: Indicadores de evaluación y criterios aplicados**

| Evaluación                       | Indicador                        | Criterio                                                                                                   |
|----------------------------------|----------------------------------|------------------------------------------------------------------------------------------------------------|
| Consistencia interna             | Alfa de Cronbach                 | Debe ser mayor o igual a 0,70.                                                                             |
|                                  | Correlación inter-ítem           | Debe ser mayor o igual a 0,30.                                                                             |
|                                  | Composite Reliability (CR)       | Debe ser mayor o igual a 0,70.                                                                             |
| Validez interna                  | Average Variance Extracted (AVE) | Debe ser mayor o igual a 0,50.                                                                             |
| Validez discriminante            | Criterio de Fornell-Larcker      | La raíz cuadrada de la AVE debe ser mayor que la correlación del constructo respectivo con cualquier otro. |
| Análisis Factorial Confirmatorio | CFI                              | Debe ser mayor o igual a 0,90.                                                                             |
|                                  | TLI                              | Debe ser mayor o igual a 0,90.                                                                             |
|                                  | RMSEA                            | Debe ser menor o igual a 0,05.                                                                             |
|                                  | Loadings / Cargas                | Deben ser mayores a 0,50 y ser estadísticamente significativos.                                            |

Inicialmente, el **análisis factorial exploratorio** reveló que la *autoeficacia* y el *control conductual percibido* debían tratarse como un único factor o constructo, dada la alta correlación entre sus ítems y su tendencia a agruparse en un solo factor.

Posteriormente, se evaluó la **consistencia interna** de cada constructo mediante el **Alfa de Cronbach** y la **correlación inter-ítem**. Basado en este análisis, se eliminaron dos ítems debido a su baja correlación inter-ítem.

A continuación, el **análisis factorial confirmatorio** llevó a la eliminación de otros cinco ítems, debido a sus bajas **cargas factoriales** (*loadings* o coeficientes). Además, se determinó que los indicadores **CFI**, **TLI** y **RMSEA** (donde los dos primeros son indicadores de bondad de ajuste y el tercero es el error cuadrático medio del modelo) se encontraban fuera de los rangos aceptables.

Finalmente, el análisis de **validez convergente y discriminante** identificó factores que no cumplían con los criterios mínimos, lo que llevó a la eliminación de 18 ítems. Esta

exclusión abarcó el factor *actitud* de la TCP (Teoría del Comportamiento Planificado), debido a sus altas cargas cruzadas con otros factores como *gravedad*, *eficacia* y *autoeficacia*. Adicionalmente, la eliminación de los ítems mencionados resultó en la desaparición del factor *Control conductual percibido*. Tras estos ajustes, se alcanzó el cumplimiento de todos los criterios estadísticos requeridos.

A continuación, se presentan en las tablas los indicadores del modelo validado para los diferentes análisis. La Tabla 3 muestra los resultados del análisis factorial exploratorio. Los valores indican una fuerte correlación de cada ítem con su factor correspondiente, mientras que la correlación con los demás factores es significativamente menor.

**Tabla 3: Análisis factorial exploratorio**

| Ítem            | VDep        | Costos      | Vulnerabilidad | Eficacia    | Normas      | Autoeficacia | Gravedad    |
|-----------------|-------------|-------------|----------------|-------------|-------------|--------------|-------------|
| vdep_1          | <b>0,83</b> | -0,03       | -0,07          | 0,00        | -0,02       | -0,02        | 0,06        |
| vdep_3          | <b>0,83</b> | 0,06        | -0,05          | 0,00        | -0,04       | 0,01         | 0,06        |
| vdep_4          | <b>0,80</b> | 0,05        | -0,07          | -0,02       | -0,01       | 0,04         | 0,00        |
| vdep_5          | <b>0,79</b> | 0,03        | 0,00           | 0,00        | -0,01       | 0,01         | 0,00        |
| vdep_2          | <b>0,78</b> | 0,00        | 0,05           | 0,03        | -0,07       | 0,04         | 0,00        |
| vdep_8          | <b>0,78</b> | -0,08       | 0,09           | -0,01       | 0,02        | -0,01        | -0,02       |
| vdep_7          | <b>0,65</b> | -0,04       | 0,06           | 0,06        | 0,12        | -0,01        | -0,06       |
| vdep_6          | <b>0,51</b> | -0,08       | 0,17           | 0,06        | 0,22        | 0,04         | -0,15       |
| costos6         | -0,01       | <b>0,82</b> | -0,06          | -0,02       | 0,03        | -0,01        | 0,00        |
| costos4         | 0,02        | <b>0,77</b> | 0,08           | 0,06        | 0,02        | -0,05        | -0,02       |
| costos2         | 0,02        | <b>0,72</b> | 0,06           | 0,06        | -0,01       | -0,06        | 0,01        |
| costos5         | -0,07       | <b>0,63</b> | 0,03           | -0,11       | -0,02       | 0,09         | -0,04       |
| costos1         | 0,02        | <b>0,45</b> | 0,08           | -0,15       | 0,16        | 0,11         | 0,08        |
| vulnerabilidad2 | -0,01       | 0,01        | <b>0,79</b>    | -0,08       | 0,05        | 0,04         | -0,01       |
| vulnerabilidad6 | 0,04        | 0,04        | <b>0,78</b>    | 0,03        | -0,04       | -0,05        | 0,05        |
| vulnerabilidad4 | -0,05       | 0,08        | <b>0,68</b>    | 0,03        | -0,04       | -0,05        | 0,01        |
| vulnerabilidad1 | 0,01        | 0,02        | <b>0,56</b>    | -0,09       | 0,19        | 0,08         | 0,06        |
| eficacia4       | -0,01       | -0,02       | -0,02          | <b>0,81</b> | 0,03        | 0,00         | 0,00        |
| eficacia5       | 0,06        | 0,00        | -0,04          | <b>0,74</b> | 0,04        | -0,02        | 0,02        |
| eficacia3       | 0,01        | 0,04        | -0,01          | <b>0,68</b> | -0,03       | 0,13         | 0,03        |
| eficacia2       | -0,01       | -0,01       | 0,02           | <b>0,56</b> | 0,04        | 0,12         | 0,09        |
| normas1         | -0,01       | 0,04        | -0,06          | -0,05       | <b>0,85</b> | 0,07         | 0,02        |
| normas2         | 0,01        | 0,01        | 0,04           | 0,04        | <b>0,80</b> | -0,06        | 0,05        |
| normas3         | -0,04       | 0,06        | 0,13           | 0,11        | <b>0,70</b> | -0,06        | -0,03       |
| autoeficacia1   | 0,09        | 0,04        | -0,05          | -0,06       | 0,00        | <b>0,74</b>  | 0,01        |
| autoeficacia2   | -0,04       | -0,02       | 0,03           | 0,12        | -0,08       | <b>0,64</b>  | 0,06        |
| autoeficacia3   | -0,07       | -0,07       | 0,10           | 0,17        | -0,10       | <b>0,56</b>  | 0,02        |
| gravedad6       | 0,04        | -0,05       | 0,01           | -0,02       | 0,05        | 0,02         | <b>0,78</b> |
| gravedad4       | -0,01       | 0,02        | 0,01           | 0,09        | -0,02       | -0,07        | <b>0,74</b> |
| gravedad2       | -0,01       | 0,03        | 0,05           | -0,03       | 0,02        | 0,14         | <b>0,60</b> |

Los resultados del **análisis factorial confirmatorio** se presentan a continuación. La Tabla 4 muestra las cargas factoriales estandarizadas y su significación estadística, evaluada mediante el valor p. Se observa que todos los coeficientes superan 0.50 y son significativos al 1% (**Tabla 4**). Además, los indicadores de bondad de ajuste, CFI y TLI, fueron de 0.94 y 0.93, respectivamente, mientras que el error cuadrático medio de aproximación, RMSEA, fue de 0.047, lo que confirma el buen ajuste del modelo.

**Tabla 4: Análisis factorial confirmatorio**

| V. Dep.        | Item            | Carga | Valor p |
|----------------|-----------------|-------|---------|
| Vulnerabilidad | vulnerabilidad1 | 0,71  | 0,00    |
|                | vulnerabilidad2 | 0,82  | 0,00    |
|                | vulnerabilidad4 | 0,69  | 0,00    |
|                | vulnerabilidad6 | 0,77  | 0,00    |
| Gravedad       | gravedad2       | 0,64  | 0,00    |
|                | gravedad4       | 0,75  | 0,00    |
|                | gravedad6       | 0,80  | 0,00    |
| Eficacia       | eficacia2       | 0,67  | 0,00    |
|                | eficacia3       | 0,74  | 0,00    |
|                | eficacia4       | 0,80  | 0,00    |
|                | eficacia5       | 0,76  | 0,00    |
| Autoeficacia   | autoeficacia1   | 0,71  | 0,00    |
|                | autoeficacia2   | 0,70  | 0,00    |
|                | autoeficacia3   | 0,65  | 0,00    |
| Costos         | costos1         | 0,56  | 0,00    |
|                | costos2         | 0,75  | 0,00    |
|                | costos4         | 0,83  | 0,00    |
|                | costos5         | 0,62  | 0,00    |
|                | costos6         | 0,78  | 0,00    |
|                | normas1         | 0,79  | 0,00    |
| Normas         | normas2         | 0,87  | 0,00    |
|                | normas3         | 0,81  | 0,00    |
| Vdep           | vdep_1          | 0,84  | 0,00    |
|                | vdep_2          | 0,80  | 0,00    |
|                | vdep_3          | 0,84  | 0,00    |
|                | vdep_4          | 0,80  | 0,00    |
|                | vdep_5          | 0,79  | 0,00    |
|                | vdep_6          | 0,51  | 0,00    |
|                | vdep_7          | 0,65  | 0,00    |
|                | vdep_8          | 0,77  | 0,00    |

La **consistencia interna** se evaluó mediante el Alfa de Cronbach y Composite Reliability (CR). Todos los factores superaron el valor mínimo de 0,70, para ambos indicadores. Por otro lado, la **validez convergente** se confirmó al verificar que la Varianza Media Extraída (AVE) fuera igual o superior a 0.50 en todos los casos. Por último, se confirmó la **validez discriminante**, al demostrar que la raíz cuadrada de la AVE de cada factor era mayor que su correlación con cualquier otro factor <sup>1</sup> (Tabla 5).

**Tabla 5: Consistencia interna, validez convergente y validez discriminante**

|                | Alfa | CR   | AVE  | Vulnerabilidad | Gravedad    | Eficacia    | Autoeficacia | Costos      | Normas      | Vdep        |
|----------------|------|------|------|----------------|-------------|-------------|--------------|-------------|-------------|-------------|
| Vulnerabilidad | 0.83 | 0.84 | 0.56 | <b>0.75</b>    |             |             |              |             |             |             |
| Gravedad       | 0.77 | 0.77 | 0.53 | 0.24           | <b>0.73</b> |             |              |             |             |             |
| Eficacia       | 0.83 | 0.83 | 0.56 | -0.05          | 0.46        | <b>0.75</b> |              |             |             |             |
| Autoeficacia   | 0.80 | 0.80 | 0.50 | -0.07          | 0.37        | 0.54        | <b>0.71</b>  |             |             |             |
| Costos         | 0.83 | 0.84 | 0.52 | 0.52           | 0.02        | -0.16       | -0.29        | <b>0.72</b> |             |             |
| Normas         | 0.86 | 0.86 | 0.68 | 0.58           | 0.23        | 0.17        | 0.03         | 0.44        | <b>0.82</b> |             |
| Vdep           | 0.91 | 0.90 | 0.56 | 0.02           | 0.20        | 0.23        | 0.33         | -0.07       | 0.06        | <b>0.75</b> |

Alfa: alfa de Cronbach; CR: composite reliability; AVE: average variance extracted. En la matriz, los valores en la diagonal son la raíz cuadrada de la AVE; los valores bajo la diagonal son correlaciones.

<sup>1</sup> Se verifica comparando el valor de la diagonal de la matriz (panel derecho de la tabla 5) con los otros valores de la fila y/o columna correspondiente.

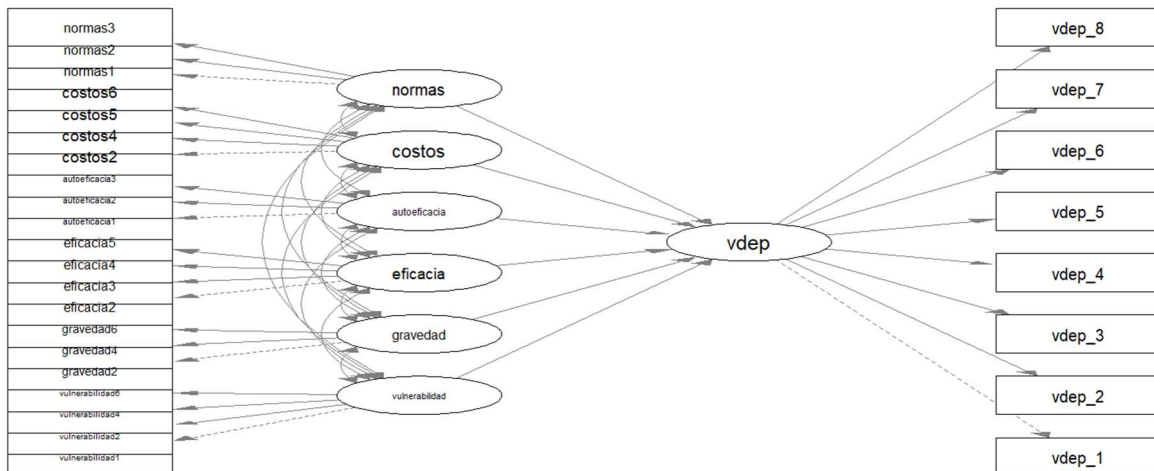
### 4.3 Modelo Estructural

La muestra utilizada para la estimación de las ecuaciones estructurales incluyó las respuestas de la Encuesta 1 y la Encuesta 2, alcanzando un total de 2.023 respuestas válidas (participantes que cumplieran los criterios de elegibilidad).

A continuación, se presentan los resultados de las regresiones, donde la variable dependiente es la motivación de protección (intención de protegerse), y las variables independientes corresponden a los componentes de la Teoría de la Motivación de Protección (TMP) y la Teoría del Comportamiento Planificado (TCP). Todas estas variables fueron modeladas como factores, es decir, combinaciones lineales de los ítems, de acuerdo con el modelo de medición desarrollado previamente.

A continuación, se describe el Modelo de Ecuación Estructural a estimar a través del modelo SEM (**Figura 1**):

**Figura 1: Modelo de Ecuación Estructural (SEM)**



Fuente: Elaboración propia.

#### Hipótesis causales

- H1: La percepción de mayor vulnerabilidad incrementa la motivación para seguir las recomendaciones de seguridad contra el fraude informático.
- H2: La percepción de mayor gravedad incrementa la motivación para seguir las recomendaciones de seguridad contra el fraude informático.
- H3: La percepción de mayor eficacia de la respuesta incrementa la motivación para seguir las recomendaciones de seguridad contra el fraude informático.
- H4: La percepción de mayor autoeficacia incrementa la motivación para seguir las recomendaciones de seguridad contra el fraude informático.
- H5: La percepción de mayores costos de respuesta reduce la motivación para seguir las recomendaciones de seguridad contra el fraude informático.
- H6: Las normas descriptivas incrementan la motivación para seguir las recomendaciones de seguridad contra el fraude informático.

Los resultados de las regresiones se presentan en la Tabla 6 (se muestran coeficientes estandarizados, permitiendo la comparación entre ellos<sup>2</sup>). La ecuación de la columna 1 incluye únicamente variables derivadas de la Teoría de la Motivación de Protección (TMP), mientras que la columna 2 incorpora adicionalmente las normas sociales provenientes de la Teoría del Comportamiento Planeado (TCP). En las columnas 3 y 4 se añaden otras variables de control, tales como características socioeconómicas, información y experiencia previa. Los resultados sugieren que la *autoeficacia* es la variable que ejerce mayor influencia sobre la motivación de protección, debido tanto a la magnitud de sus coeficientes como a su elevada significancia estadística en todas las ecuaciones estimadas. La *gravedad percibida* aparece como la segunda variable más relevante, mientras que la *eficacia de la respuesta* y los *costos de la respuesta* presentan resultados menos robustos, ya que pierden significancia estadística al incluir variables adicionales. Cabe destacar que, en general, los coeficientes obtenidos se encuentran en un rango similar al reportado en investigaciones previas (Ifinedo, 2012; Jansen & Van Schaik, 2018b) (**Tabla 6**).

**Tabla 6: Resultados de ecuación estructural**

|                | V.Dep: Motivación de protección |                     |                     |                     |
|----------------|---------------------------------|---------------------|---------------------|---------------------|
|                | (1)                             | (2)                 | (3)                 | (4)                 |
| Vulnerabilidad | 0,025<br>(0,032)                | 0,011<br>(0,037)    | -0,013<br>(0,037)   | -0,019<br>(0,037)   |
| Gravedad       | 0,090**<br>(0,040)              | 0,091**<br>(0,040)  | 0,100**<br>(0,039)  | 0,096**<br>(0,039)  |
| Eficacia       | 0,086**<br>(0,036)              | 0,079**<br>(0,038)  | 0,060<br>(0,038)    | 0,058<br>(0,038)    |
| Autoeficacia   | 0,186***<br>(0,041)             | 0,187***<br>(0,041) | 0,205***<br>(0,042) | 0,222***<br>(0,041) |
| Costos         | -0,065**<br>(0,033)             | -0,072**<br>(0,034) | -0,049<br>(0,034)   | -0,051<br>(0,034)   |
| Normas         |                                 | 0,029<br>(0,035)    | 0,012<br>(0,035)    | 0,012<br>(0,035)    |
| Controles      |                                 |                     | 1                   | 2                   |
| N              | 2.023                           | 2.023               | 2.023               | 2.023               |
| CFI            | 0,964                           | 0,961               | 0,942               | 0,927               |
| TLI            | 0,958                           | 0,956               | 0,937               | 0,920               |
| RMSEA          | 0,038                           | 0,038               | 0,033               | 0,036               |

Se muestran coeficientes estandarizados (comparables).

Entre paréntesis se muestran los errores estándar robustos de cada estimación.

Controles: 1 = dummies de género, edad, educación e ingreso;

2 = controles 1 + dummies de información y experiencia.

\*p<0,1; \*\*p<0,05; \*\*\*p<0,01

En cuanto a las variables de control utilizadas, destaca el efecto negativo y estadísticamente significativo asociado al género femenino, con niveles de significancia del 10% en la ecuación 3 y del 5% en la ecuación 4. Este hallazgo sugiere que la

<sup>2</sup> los coeficientes estandarizados son utilizados especialmente cuando las variables observadas no comparten una escala común o tienen escalas arbitrarias. Su propósito principal es facilitar la interpretación de la magnitud y la dirección de las relaciones entre las variables al colocarlas en una escala comparable.

motivación de protección es menor en las mujeres, incluso después de controlar por las variables derivadas de la TMP, TCP y otras características socioeconómicas.

#### 4.4 Análisis de Robustez

Para evaluar la robustez de los resultados, se repitieron las estimaciones utilizando mínimos cuadrados ordinarios (MCO). En estas regresiones, tanto las variables dependientes como independientes corresponden a las predicciones del modelo de medición. La columna 2 presenta los resultados de una especificación análoga a la columna 2 de la tabla anterior, confirmando la significancia estadística y aumentando la magnitud del efecto de la *autoeficacia*. Por su parte, la columna 1 muestra que, al excluir la autoeficacia de la ecuación, el  $R^2$  disminuye en un 21% (-2,7 puntos porcentuales), reafirmando así su relevancia como determinante de la motivación de protección. Las columnas 3 y 4 reproducen, respectivamente, las especificaciones de las columnas 3 y 4 anteriores, mientras que en las columnas 5 y 6 se incluyen términos cuadráticos e interacciones. En todas estas estimaciones, la *autoeficacia* mantiene el coeficiente más alto en términos de magnitud y significancia estadística. La *gravedad percibida* y los *costos de la respuesta* también resultan significativos en la mayoría de las regresiones realizadas (**Tabla 7**).

**Tabla 7: Resultados de estimaciones MCO**

|                     | V.Dep: Motivación de protección |                     |                     |                     |                     |                     |
|---------------------|---------------------------------|---------------------|---------------------|---------------------|---------------------|---------------------|
|                     | (1)                             | (2)                 | (3)                 | (4)                 | (5)                 | (6)                 |
| Vulnerabilidad      | 0,006<br>(0,033)                | 0,009<br>(0,032)    | -0,012<br>(0,033)   | -0,019<br>(0,033)   | 0,029<br>(0,037)    | 0,009<br>(0,033)    |
| Gravedad            | 0,166***<br>(0,034)             | 0,089***<br>(0,033) | 0,098***<br>(0,033) | 0,093***<br>(0,033) | 0,060<br>(0,054)    | 0,104**<br>(0,041)  |
| Eficacia            | 0,146***<br>(0,030)             | 0,051<br>(0,032)    | 0,033<br>(0,033)    | 0,032<br>(0,033)    | 0,005<br>(0,045)    | 0,048<br>(0,033)    |
| Autoeficacia        |                                 | 0,234***<br>(0,034) | 0,250***<br>(0,034) | 0,273***<br>(0,036) | 0,302***<br>(0,046) | 0,236***<br>(0,036) |
| Costos              | -0,134***<br>(0,028)            | -0,067**<br>(0,030) | -0,047<br>(0,030)   | -0,050*<br>(0,030)  | -0,056*<br>(0,030)  | -0,062**<br>(0,030) |
| Normas              | 0,033<br>(0,031)                | 0,036<br>(0,031)    | 0,020<br>(0,031)    | 0,019<br>(0,031)    | 0,041<br>(0,035)    | 0,036<br>(0,031)    |
| Controles           |                                 |                     | 1                   | 2                   | 3                   | 4                   |
| N                   | 2.023                           | 2.023               | 2.023               | 2.023               | 2.023               | 2.023               |
| R <sup>2</sup>      | 0,097                           | 0,124               | 0,136               | 0,140               | 0,128               | 0,124               |
| Adj. R <sup>2</sup> | 0,095                           | 0,121               | 0,128               | 0,131               | 0,122               | 0,121               |

Se muestran coeficientes estandarizados (comparables).

Entre paréntesis se muestran los errores estándar robustos de cada estimación.

Controles: 1 = dummies de género, edad, educación e ingreso;

2 = controles 1 + dummies de información y experiencia.

3 = efectos cuadráticos. 4 = términos de interacción.  
\*p<0,1; \*\*p<0,05; \*\*\*p<0,01

Finalmente, los gráficos de dispersión entre la variable dependiente y cada una de las variables independientes revelaron la presencia de numerosos valores atípicamente bajos (outliers), que podrían estar atenuando las estimaciones. Para evaluar esta hipótesis se realizaron regresiones por cuantiles<sup>3</sup>, las cuales son menos sensibles a la presencia de valores extremos. Las columnas presentan estimaciones para los percentiles 0,25; 0,50 y 0,75, respectivamente. La columna 2, que corresponde a la regresión mediana (percentil 0,50), es la más comparable con las estimaciones anteriores. En todas las estimaciones por cuantiles aumentó la magnitud de los coeficientes respecto a resultados previos, salvo en el caso de la *gravedad percibida*. Se observa nuevamente que la *autoeficacia* mantiene significancia estadística en todas las regresiones. Sin embargo, la *gravedad percibida* pierde significancia, mientras que la *vulnerabilidad percibida* y las *normas subjetivas*—que previamente no habían resultado significativas—adquieren ahora significancia estadística (**Tabla 8**).

**Tabla 8: Resultados de regresiones de cuantiles**

|                | V.Dep: Motivación de protección |                     |                      |
|----------------|---------------------------------|---------------------|----------------------|
|                | 0,25<br>(1)                     | 0,50<br>(2)         | 0,75<br>(3)          |
| Vulnerabilidad | 0,051***<br>(0,017)             | 0,036<br>(0,027)    | 0,019**<br>(0,007)   |
| Gravedad       | 0,042<br>(0,028)                | 0,053<br>(0,044)    | 0,012<br>(0,021)     |
| Eficacia       | 0,117***<br>(0,027)             | 0,170***<br>(0,028) | 0,068***<br>(0,016)  |
| Autoeficacia   | 0,241***<br>(0,032)             | 0,314***<br>(0,044) | 0,126***<br>(0,019)  |
| Costos         | -0,021<br>(0,013)               | -0,084**<br>(0,034) | -0,011***<br>(0,004) |
| Normas         | 0,042***<br>(0,015)             | 0,073***<br>(0,026) | 0,023***<br>(0,007)  |
| N              | 2.023                           | 2.023               | 2.023                |

Se muestran coeficientes estandarizados (comparables).

Entre paréntesis se muestran los errores estándar robustos de cada estimación.

\*p<0,1; \*\*p<0,05; \*\*\*p<0,01

<sup>3</sup> A diferencia de los modelos clásicos de regresión, la regresión por cuantiles (*quantile regression*, QR) permite analizar el efecto de las variables predictoras en diferentes puntos (cuantiles) de la distribución condicional de la variable dependiente, en lugar de enfocarse únicamente en la media condicional, como ocurre con los mínimos cuadrados ordinarios (OLS). Este método facilita identificar y modelar efectos heterogéneos de las variables independientes sobre la variable respuesta. Además, la regresión por cuantiles es menos sensible a la influencia de valores atípicos (*outliers*) en la variable dependiente en comparación con la regresión OLS, ya que sus coeficientes se estiman minimizando una suma ponderada de residuos absolutos en lugar de residuos elevados al cuadrado (Davino et al., 2013).

En conclusión, este análisis confirma que la *autoeficacia* es la variable clave para explicar la *motivación de protección*, destacándose por su efecto robusto, altamente significativo en todas las especificaciones funcionales evaluadas y con una magnitud considerablemente superior respecto a otros factores analizados. En segundo lugar, también se destacan la *gravedad percibida* y los *costos de la respuesta*; sin embargo, las estimaciones para estas variables resultaron menos consistentes y robustas.

## 4.5 Análisis de Heterogeneidad

En esta sección se presenta un análisis de heterogeneidad de los resultados mediante estimaciones por subgrupos de la muestra. Este enfoque permite comprender cómo varían los efectos principales según diferentes segmentos. Para ello, la muestra fue dividida en subgrupos relativamente homogéneos, utilizando las variables y criterios que se detallan en la **Tabla 9**.

**Tabla 9: Subgrupos**

| Variable      | Subgrupo 1                 | N1    | Subgrupo 2                          | N2    |
|---------------|----------------------------|-------|-------------------------------------|-------|
| Género        | Femenino                   | 945   | Masculino                           | 1.051 |
| Edad          | Hasta 44 años (1, 2)       | 834   | 45 años o más (3, 4)                | 1.189 |
| Educación     | Hasta E. Técnica (1, 2, 3) | 870   | E. Universitaria y postgrado (4, 5) | 1.153 |
| Ingreso       | Hasta \$900.000 (1, 2, 3)  | 937   | \$901.000 o más (4, 5, 6)           | 1.086 |
| Informado*    | Poco informado (1, 2, 3)   | 1.178 | Bien informado (4, 5)               | 845   |
| Experiencia** | Poca experiencia (1, 2, 3) | 1.024 | Alta experiencia (4, 5)             | 999   |
| Fraude***     | Ha sufrido fraude          | 622   | No ha sufrido fraude                | 1.401 |

Entre paréntesis se muestran las categorías utilizadas en las variables ordinales.

\*¿Qué tan bien informado(a) se siente sobre los riesgos del fraude informático?

\*\*¿Se considera experimentado(a) en el uso del sitio web o aplicaciones de su institución financiera?

\*\*\* ¿Usted ha sido víctima de fraude informático en la banca en línea durante el último año?

Cabe destacar que, en este análisis, los valores-p asociados a las pruebas de hipótesis se corrigieron utilizando el método de Benjamini y Hochberg (1995), con el propósito de controlar la tasa de errores tipo I (falsos positivos) al realizar múltiples comparaciones. La tabla presenta los resultados del análisis por subgrupos según género y edad. Nuevamente se confirma que la *autoeficacia* es la variable más robusta, manteniéndose significativa en todos los subgrupos analizados. Por otro lado, el análisis segmentado por género revela que la *gravedad percibida* es particularmente relevante para explicar la motivación de protección en mujeres, mientras que la *eficacia de la respuesta* resulta más influyente en hombres (**Tabla 12**).

**Tabla 12: Resultados para subgrupos: Género y Edad**

V.Dep: Motivación de protección

|                | Femenino           | Masculino        | Edad: -44        | Edad: 45+         |
|----------------|--------------------|------------------|------------------|-------------------|
|                | (1)                | (2)              | (3)              | (4)               |
| Vulnerabilidad | -0,065<br>(0,050)  | 0,078<br>(0,054) | 0,051<br>(0,066) | -0,047<br>(0,045) |
| Gravedad       | 0,179**<br>(0,055) | 0,008<br>(0,052) | 0,050<br>(0,062) | 0,114<br>(0,053)  |
| Eficacia       | 0,030              | 0,125*           | 0,065            | 0,071             |

|              |         |          |         |          |
|--------------|---------|----------|---------|----------|
|              | (0,054) | (0,055)  | (0,057) | (0,052)  |
| Autoeficacia | 0,175** | 0,212*** | 0,201** | 0,202*** |
|              | (0,057) | (0,062)  | (0,071) | (0,053)  |
| Costos       | -0,074  | -0,058   | -0,105  | -0,030   |
|              | (0,050) | (0,048)  | (0,051) | (0,048)  |
| Normas       | 0,083   | -0,039   | -0,009  | 0,043    |
|              | (0,049) | (0,047)  | (0,056) | (0,044)  |
| N            | 945     | 1.051    | 834     | 1.189    |

Se muestran coeficientes estandarizados (comparables).

Entre paréntesis se muestran los errores estándar robustos de cada estimación.

\*p<0,1; \*\*p<0,05; \*\*\*p<0,01

La Tabla 13 presenta los resultados correspondientes a los subgrupos definidos según nivel educacional e ingreso. Se confirma nuevamente un efecto robusto y consistente de la *autoeficacia*. En el subgrupo con menor nivel educacional (hasta educación superior técnica), los *costos de la respuesta* exhiben un efecto estadísticamente significativo. En cuanto al análisis por nivel de ingreso, la gravedad percibida alcanza significancia al 10% en el grupo de menores ingresos, mientras que en el subgrupo de mayores ingresos destaca la *eficacia de la respuesta*, con un efecto significativo y de elevada magnitud (Tabla 13).

**Tabla 13: Resultados para subgrupos: Educación e Ingreso**

V.Dep: Motivación de protección

|                | -Ed.Técnica | Ed.Universitaria+ | Ingreso:-900M | Ingreso:901M+ |
|----------------|-------------|-------------------|---------------|---------------|
|                | (1)         | (2)               | (3)           | (4)           |
| Vulnerabilidad | -0,045      | 0,060             | -0,003        | 0,027         |
|                | (0,057)     | (0,050)           | (0,055)       | (0,052)       |
| Gravedad       | 0,106       | 0,041             | 0,156*        | 0,043         |
|                | (0,054)     | (0,060)           | (0,062)       | (0,055)       |
| Eficacia       | 0,085       | 0,082             | -0,015        | 0,135**       |
|                | (0,062)     | (0,047)           | (0,058)       | (0,049)       |
| Autoeficacia   | 0,155*      | 0,226***          | 0,225***      | 0,159**       |
|                | (0,063)     | (0,055)           | (0,063)       | (0,054)       |
| Costos         | -0,133**    | -0,013            | -0,088        | -0,067        |
|                | (0,050)     | (0,047)           | (0,050)       | (0,048)       |
| Normas         | 0,129       | -0,040            | 0,088         | -0,017        |
|                | (0,058)     | (0,044)           | (0,050)       | (0,050)       |
| N              | 870         | 1.153             | 937           | 1.086         |

Se muestran coeficientes estandarizados (comparables).

Entre paréntesis se muestran los errores estándar robustos de cada estimación.

\*p<0,1; \*\*p<0,05; \*\*\*p<0,01

Finalmente, se presentan los resultados para los subgrupos definidos según las variables '*informado*' ("¿Qué tan bien informado(a) se siente sobre los riesgos del fraude informático?"), '*experiencia*' ("¿Se considera experimentado(a) en el uso del sitio web o aplicaciones de su institución financiera?") y '*fraude*' ("¿Ha sido víctima de fraude informático en la banca en línea durante el último año?"). La columna 1 muestra el único caso en que la *autoeficacia* no resulta significativa. Por otra parte, se identifican efectos significativos de la *gravedad percibida* sobre la *motivación de protección* en personas menos informadas; de la *eficacia de la respuesta* en aquellas personas que no han experimentado fraude; y de los *costos de la respuesta* en usuarios menos informados y experimentados en el uso de la banca en línea (**Tabla 14**).

**Tabla 14: Resultados para subgrupos: Información, experiencia y fraude**

V.Dep: Motivación de protección

|                | -Informado         | +Informado          | -Experiencia        | +Experiencia        | No Fraude          | Si Fraude          |
|----------------|--------------------|---------------------|---------------------|---------------------|--------------------|--------------------|
|                | (1)                | (2)                 | (3)                 | (4)                 | (5)                | (6)                |
| Vulnerabilidad | 0,014<br>(0,061)   | 0,011<br>(0,047)    | 0,013<br>(0,057)    | 0,008<br>(0,049)    | 0,021<br>(0,043)   | -0,021<br>(0,067)  |
| Gravedad       | 0,160**<br>(0,060) | -0,010<br>(0,060)   | 0,124<br>(0,060)    | 0,001<br>(0,050)    | 0,095<br>(0,053)   | 0,070<br>(0,068)   |
| Eficacia       | 0,091<br>(0,057)   | 0,059<br>(0,050)    | 0,099<br>(0,061)    | 0,041<br>(0,047)    | 0,119*<br>(0,052)  | 0,032<br>(0,060)   |
| Autoeficacia   | 0,112<br>(0,056)   | 0,286***<br>(0,067) | 0,151**<br>(0,056)  | 0,307***<br>(0,072) | 0,161**<br>(0,052) | 0,222**<br>(0,068) |
| Costos         | -0,121*<br>(0,051) | -0,036<br>(0,044)   | -0,127**<br>(0,047) | -0,010<br>(0,049)   | -0,054<br>(0,042)  | -0,087<br>(0,060)  |
| Normas         | 0,052<br>(0,058)   | 0,028<br>(0,044)    | 0,073<br>(0,057)    | -0,014<br>(0,044)   | 0,015<br>(0,044)   | 0,054<br>(0,057)   |
| N              | 845                | 1.178               | 1.024               | 999                 | 1.401              | 622                |

Se muestran coeficientes estandarizados (comparables).

Entre paréntesis se muestran los errores estándar robustos de cada estimación.

\*p<0,1; \*\*p<0,05; \*\*\*p<0,01

En conclusión, este análisis ratifica que la *autoeficacia* es la variable con mayor efecto sobre la motivación de protección, destacándose tanto por la magnitud de su coeficiente como por su alta significancia y robustez en diversas especificaciones. En segundo y tercer lugar se encuentran la *gravedad percibida* y los *costos de la respuesta*, respectivamente, que también constituyen factores relevantes para explicar la motivación de protección, aunque con efectos menores en magnitud y menos consistentes en términos de robustez.

## 5. Conclusiones

El presente estudio ha revelado importantes hallazgos sobre el comportamiento precautorio de los consumidores ante el riesgo de fraude en los medios de pago en línea en Chile. Mediante el análisis de una encuesta aplicada a más de 2.000 consumidores, se identificaron factores clave que influyen en la motivación de los consumidores para adoptar medidas de protección frente al fraude informático. De estos factores, la *autoeficacia* destaca como el más significativo y robusto, indicando que la creencia de los individuos en su propia capacidad para implementar recomendaciones de seguridad es fundamental para motivar comportamientos protectores. Este resultado subraya la importancia de empoderar a los consumidores mediante estrategias educativas y campañas dirigidas a fortalecer su confianza en su habilidad para protegerse eficazmente.

Asimismo, la *percepción sobre la gravedad del fraude*, aunque menos robusta, también desempeña un papel relevante. Los consumidores que perciben mayores consecuencias negativas ante la posibilidad de ser víctimas están más motivados para implementar medidas preventivas. Este hallazgo sugiere que una comunicación clara y efectiva acerca de los riesgos y las consecuencias asociadas al fraude puede constituir una herramienta valiosa para promover comportamientos de protección.

Finalmente, los *costos percibidos de la respuesta* (en términos de tiempo, dinero y esfuerzo) afectan negativamente la motivación para adoptar comportamientos seguros. Por lo tanto, reducir estos costos percibidos mediante la simplificación y accesibilidad de las medidas de seguridad podría incrementar significativamente la disposición de los consumidores a seguir las recomendaciones de protección.

El análisis por subgrupos reveló diferencias significativas en la motivación de protección según el género. Particularmente, las mujeres mostraron una menor motivación protectora, incluso después de controlar por otros factores relevantes, y además reportaron haber sido víctimas de fraude con mayor frecuencia que los hombres. Esto pone en evidencia la importancia de desarrollar estrategias diferenciadas y personalizadas en las campañas de sensibilización y educación, tomando en cuenta las características y necesidades específicas de cada segmento demográfico.

La metodología aplicada, basada en modelos de ecuaciones estructurales, permitió una comprensión profunda y detallada de las relaciones entre las variables estudiadas. La robustez de los resultados fue corroborada mediante múltiples pruebas de validación y análisis de heterogeneidad, lo que refuerza la confianza en las conclusiones obtenidas. Desde la perspectiva de las políticas públicas, los resultados ofrecen evidencia empírica sólida para el diseño de intervenciones eficaces en la prevención del fraude informático.

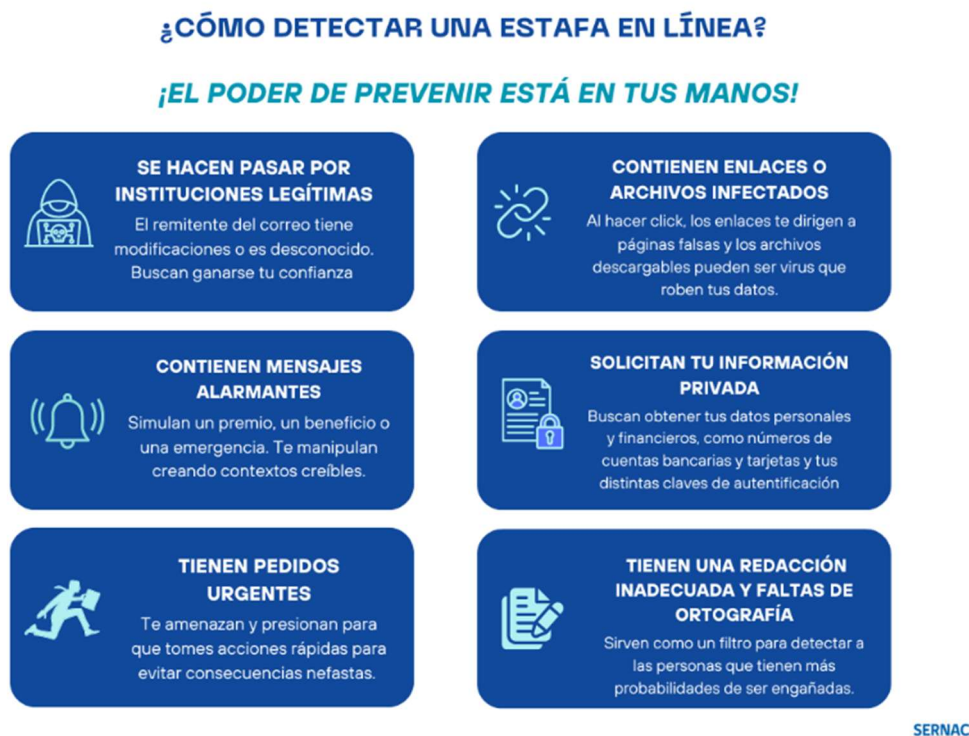
Las campañas deben centrarse especialmente en fortalecer la *autoeficacia* de los consumidores, destacando mensajes motivacionales que empoderen a los usuarios y refuercen su capacidad para protegerse. Además, resulta esencial comunicar claramente las graves consecuencias que implica ser víctima de fraude, y facilitar la adopción de

medidas de seguridad mediante la reducción de los costos percibidos por los consumidores.

Finalmente, se recomienda que futuras investigaciones profundicen en la identificación de mecanismos específicos capaces de fortalecer la autoeficacia de los consumidores y que evalúen el impacto de estas intervenciones mediante diseños experimentales. La aplicación de métodos experimentales no solo incrementaría la validez interna de los hallazgos, sino que también facilitaría la comprensión de los mecanismos subyacentes, contribuyendo así a un conocimiento más profundo sobre cómo los consumidores adoptan medidas de protección. Por ejemplo, aunque se reconoce que apelar a la autoeficacia incentiva la motivación de protección, aún no está claro cuál es la manera más efectiva de hacerlo. Campañas desarrolladas bajo este enfoque podrían ser implementadas por el SERNAC y las instituciones financieras, con el fin de motivar a los consumidores a protegerse activamente contra el fraude informático.

Como ilustración, en la Figura 3 se presenta una propuesta de campaña que recomienda acciones concretas para detectar correos de phishing, acompañadas por un mensaje motivacional orientado a fortalecer la autoeficacia de los consumidores: “¡El poder de prevenir está en tus manos!”. Este tipo de mensajes busca empoderar a los usuarios, reforzando la percepción de que poseen las capacidades necesarias para protegerse, aumentando así la probabilidad de que adopten las recomendaciones sugeridas (**Figura 2**).

**Figura 2: Campaña de ciberseguridad incorporando mensaje motivacional**



## 6. Bibliografía

- Acquisti, Alessandro, Laura Brandimarte, y George Loewenstein. (2015). «Privacy and human behavior in the age of information». *Science* 347 (6221): 509-14. <https://doi.org/10.1126/science.aaa1465>.
- Ajzen, Icek. (1991). «The theory of planned behavior». *Organizational Behavior and Human Decision Processes, Theories of Cognitive Self-Regulation*, 50 (2): 179-211. [https://doi.org/10.1016/0749-5978\(91\)90020-T](https://doi.org/10.1016/0749-5978(91)90020-T).
- Ajzen, Icek. (2002). «Perceived Behavioral Control, Self-Efficacy, Locus of Control, and the Theory of Planned Behavior1». *Journal of Applied Social Psychology* 32 (4): 665-83. <https://doi.org/10.1111/j.1559-1816.2002.tb00236.x>.
- Anderson, Catherine L., y Ritu Agarwal. (2010). «Practicing safe computing: a multimedia empirical examination of home computer user security behavioral intentions». *MIS Quarterly* 34 (3): 613-43.
- Bandura, Albert. (1977). «Self-efficacy: Toward a unifying theory of behavioral change». *Psychological Review* 84 (2): 191-215. <https://doi.org/10.1037/0033-295X.84.2.191>.
- Beaujean, A. A. (2014). *Latent variable modeling using R: A step-by-step guide*. Routledge.
- Benjamini, Yoav, y Yosef Hochberg. (1995). «Controlling the False Discovery Rate: A Practical and Powerful Approach to Multiple Testing». *Journal of the Royal Statistical Society: Series B (Methodological)* 57 (1): 289-300. <https://doi.org/10.1111/j.2517-6161.1995.tb02031.x>.
- Brown, T. A. (2015). *Confirmatory factor analysis for applied research* (2nd ed.). The Guilford Press.
- Brown, J. Owen, Jerry Hays, y Martin T. Stuebs Jr. (2016). «Modeling Accountant Whistleblowing Intentions: Applying the Theory of Planned Behavior and the Fraud Triangle». *Accounting and the Public Interest* 16 (1): 28-56. <https://doi.org/10.2308/apin-51675>.
- Browne, M. W., & Cudeck, R. (1993). Alternative ways of assessing model fit. En K. A. Bollen & J. S. Long (Eds.), *Testing structural equation models* (pp. 136-162). Newbury Park, CA: Sage
- Burns, Sarah, y Lynne Roberts. (2013). «Applying the Theory of Planned Behaviour to Predicting Online Safety Behaviour». *Crime Prevention and Community Safety* 15 (1): 48-64. <https://doi.org/10.1057/cpcs.2012.13>.
- Chen, Xiaofeng, Liqiang Chen, y Dazhong Wu. (2018). «Factors That Influence Employees' Security Policy Compliance: An Awareness-Motivation-Capability Perspective». *Journal of Computer Information Systems* 58 (4): 312-24. <https://doi.org/10.1080/08874417.2016.1258679>.
- Clark, L. A. & Watson, D. (1995). Constructing validity: Basic issues in objective scale development. *Psychological Assessment* 7(3): 309-319.
- Comrey, A. L., & Lee, H. B. (1992). *A first course in factor analysis* (2nd ed.). Hillsdale, NJ: Lawrence Erlbaum Associates.
- Crossler, Robert E. (2010). «Protection Motivation Theory: Understanding Determinants to Backing Up Personal Data». En *2010 43rd Hawaii International Conference on System Sciences*, 1-10. <https://doi.org/10.1109/HICSS.2010.311>.

- Crossler, Robert, y France Bélanger. (2014). «An Extended Perspective on Individual Security Behaviors: Protection Motivation Theory and a Unified Security Practices (USP) Instrument». *ACM SIGMIS Database: the DATABASE for Advances in Information Systems* 45 (4): 51-71. <https://doi.org/10.1145/2691517.2691521>.
- Davino, C., Furno, M., & Vistocco, D. (2013). *Quantile regression: Theory and applications*. Wiley.
- DeVellis, R. F. (2016). *Scale development: Theory and applications* (4th ed.). Sage Publications.
- Everitt, B. S. and Skrondal, A. (2010). *The Cambridge Dictionary of Statistics*. Cambridge University Press, 4th edition.
- Fabrigar, L. R., & Wegener, D. T. (2012). *Exploratory factor analysis*. Oxford University Press.
- Fan, Yi, Jiquan Chen, Gabriela Shirkey, Ranjeet John, Susie R. Wu, Hogeun Park, y Changliang Shao. (2016). «Applications of structural equation modeling (SEM) in ecological studies: an updated review». *Ecological Processes* 5 (1): 19. <https://doi.org/10.1186/s13717-016-0063-3>.
- Fishbein, Martin, y Icek Ajzen. (2010). *Predicting and changing behavior: The reasoned action approach*. Predicting and changing behavior: The reasoned action approach. New York, NY, US: Psychology Press.
- Floyd, Donna L., Steven Prentice-Dunn, y Ronald W. Rogers. (2000). «A meta-analysis of research on protection motivation theory». *Journal of Applied Social Psychology* 30 (2): 407-29. <https://doi.org/10.1111/j.1559-1816.2000.tb02323.x>.
- Furnell, S. M., A. Jusoh, y D. Katsabas. (2006). «The challenges of understanding and using security: A survey of end-users». *Computers & Security* 25 (1): 27-35. <https://doi.org/10.1016/j.cose.2005.12.004>.
- George, D., & Mallery, P. (2003). *SPSS for Windows step by step: A simple guide and reference* (4th ed.). Allyn & Bacon.
- George, Joey F. (2004). «The theory of planned behavior and Internet purchasing». *Internet Research* 14 (3): 198-212. <https://doi.org/10.1108/10662240410542634>.
- Hair, J. F., Hult, G. T. M., Ringle, C. M., & Sarstedt, M. (2017). *A primer on partial least squares structural equation modeling (PLS-SEM)* (2nd ed.). SAGE Publications.
- Harrington, D. (2008). *Confirmatory factor analysis (Pocket guides to social work research methods)*. Oxford University Press.
- Hoyle, R. H. (Ed.). (2012). *Handbook of structural equation modeling*. The Guilford Press.
- Hoyle, R. H. (Ed.). (2023). *Handbook of structural equation modeling* (2ª ed.). Guilford Publications.
- Hu, L., & Bentler, P. M. (1999). Cutoff criteria for fit indexes in covariance structure analysis: Conventional criteria versus new alternatives. *Structural Equation Modeling: A Multidisciplinary Journal*, 6(1), 1-55.
- Huang, Ding-Long, Pei-Luen Patrick Rau, y Gavriel Salvendy. (2010). «Perception of information security». *Behaviour & Information Technology* 29 (3): 221-32. <https://doi.org/10.1080/01449290701679361>.
- Ifinedo, Princely. (2012). «Understanding information systems security policy compliance: An integration of the theory of planned behavior and the protection motivation theory». *Computers & Security* 31 (1): 83-95. <https://doi.org/10.1016/j.cose.2011.10.007>.

- Jansen, Jurjen, y Paul van Schaik. (2018b). «Testing a model of precautionary online behaviour: The case of online banking». *Computers in Human Behavior* 87 (octubre): 371-83. <https://doi.org/10.1016/j.chb.2018.05.010>.
- Jansen, Jurjen, y Rutger Leukfeldt. (2015). «How people help fraudsters steal their money: an analysis of 600 online banking fraud cases». En *2015 Workshop on Socio-Technical Aspects in Security and Trust*, 24-31. <https://doi.org/10.1109/STAST.2015.12>.
- Johnson, R. L., & Morgan, G. B. (2016). *Survey scales: A guide to development, analysis, and reporting*. Guilford Press.
- Kaplan, D. W. (2008). *Structural equation modeling: Foundations and extensions* (2ª ed.). SAGE Publications.
- Kline, R. B. (2016). *Principles and practice of structural equation modeling* (4th ed.). The Guilford Press.
- Lee, Younghwa. (2011). «Understanding anti-plagiarism software adoption: An extended protection motivation theory perspective». *Decision Support Systems* 50 (2): 361-69. <https://doi.org/10.1016/j.dss.2010.07.009>.
- Liang, Huigang, y Yajiong, Xue. (2010). «Understanding Security Behaviors in Personal Computer Usage: A Threat Avoidance Perspective». *Journal of the Association for Information Systems* 11 (7). <https://doi.org/10.17705/1jais.00232>.
- Maddux, James E., y Ronald W. Rogers. (1983). «Protection motivation and self-efficacy: A revised theory of fear appeals and attitude change». *Journal of Experimental Social Psychology* 19 (5): 469-79. [https://doi.org/10.1016/0022-1031\(83\)90023-9](https://doi.org/10.1016/0022-1031(83)90023-9).
- Maruyama, G. M. (1998). *Basics of structural equation modeling*. Sage Publications.
- Mayhew, Matthew J., Steven M. Hubbard, Cynthia J. Finelli, Trevor S. Harding, y Donald D. Carpenter. (2009). «Using Structural Equation Modeling to Validate the Theory of Planned Behavior as a Model for Predicting Student Cheating». *The Review of Higher Education* 32 (4): 441-68.
- Mehmetoglu, M., & Venturini, S. (2021). *Structural equation modelling with partial least squares using Stata and R: Theory and applications using Stata and R*. Sage Publications.
- Mehmetoglu, M., & Venturini, S. (2022). *Structural equation modelling with partial least squares using Stata and R: Theory and applications*. CRC Press.
- Mehmetoglu, M., & Venturini, S. (2021). *Structural equation modelling with partial least squares using Stata and R: Theory and applications using Stata and R*. Sage Publications.
- Menard, Philip, Gregory J. Bott, y Robert E. Crossler. (2017). «User Motivations in Protecting Information Security: Protection Motivation Theory Versus Self-Determination Theory». *Journal of Management Information Systems* 34 (4): 1203-30. <https://doi.org/10.1080/07421222.2017.1394083>.
- Milne, Sarah, Paschal Sheeran, y Sheina Orbell. (2000). «Prediction and intervention in health-related behavior: A meta-analytic review of protection motivation theory». *Journal of Applied Social Psychology* 30 (1): 106-43. <https://doi.org/10.1111/j.1559-1816.2000.tb02308.x>.
- Moore, Tyler, y Ross Anderson. (2011). «Economics and Internet Security: A Survey of Recent Analytical, Empirical, and Behavioral Research». <https://dash.harvard.edu/handle/1/23574266>.

- Netemeyer, R. G., Bearden, W. O., & Sharma, S. (2003). *Scaling procedures: Issues and applications*. SAGE Publications.
- Ng Boon-Yuen, Atreyi Kankanhalli, y Yunjie (Calvin) Xu. (2009). «Studying users' computer security behavior: A health belief perspective». *Decision Support Systems, IT Decisions in Organizations*, 46 (4): 815-25. <https://doi.org/10.1016/j.dss.2008.11.010>.
- O'Keefe, Daniel James. (2016). *Persuasion: Theory and Research*. Thousand Oaks, CA: SAGE Publications, Inc.
- Ophoff, Jacques, y Mcguigan Lakay. (2019). «Mitigating the ransomware threat: Information Security South Africa». Editado por Hein Venter, Marianne Looek, Marijke Coetzee, Mariki Eloff, y Jan Eloff. *Information security, Communications in Computer and Information Science (CCIS)*, enero, 163-75. [https://doi.org/10.1007/978-3-030-11407-7\\_12](https://doi.org/10.1007/978-3-030-11407-7_12).
- O'Rourke, N., & Hatcher, L. (2013). *A step-by-step approach to using SAS for factor analysis and structural equation modeling* (2ª ed.). SAS Institute.
- Pett, M. A., Lackey, N. R., & Sullivan, J. J. (2003). *Making sense of factor analysis: The use of factor analysis for instrument development in health care research*. SAGE Publications.
- Ramlall, I. (2016). *Applied structural equation modelling for researchers and practitioners: Using R and Stata for behavioural research*. Emerald Group Publishing.
- Raykov, T., & Marcoulides, G. A. (2006). *A first course in structural equation modeling* (2ª ed.). Lawrence Erlbaum Associates.
- Rhee, Hyeun-Suk, Cheongtag Kim, y Young U. Ryu. (2009). «Self-efficacy in information security: Its influence on end users' information security practice behavior». *Computers & Security* 28 (8): 816-26. <https://doi.org/10.1016/j.cose.2009.05.008>.
- Robinson, J. P., Shaver, P. R., & Wrightsman, L. S. (1991). Criteria for scale selection and evaluation. In J. P. Robinson, P. R. Shaver, & L. S. Wrightsman (Eds.), *Measures of personality and social psychological attitudes* (pp. 1-15). San Diego: Academic Press.
- Rogers, Ronald W. (1975). «A Protection Motivation Theory of Fear Appeals and Attitude Change1». *The Journal of Psychology* 91 (1): 93-114. <https://doi.org/10.1080/00223980.1975.9915803>.
- Schumacker, R. E., & Lomax, R. G. (2016). *A beginner's guide to structural equation modeling* (4ª ed.). Routledge.
- Shih, Ya-Yueh, y Kwoting Fang. (2004). «The use of a decomposed theory of planned behavior to study Internet banking in Taiwan». *Internet Research* 14 (3): 213-23. <https://doi.org/10.1108/10662240410542643>.
- Slaney, Kathleen L., y Timothy P. Racine. (2013). «What's in a name? Psychology's ever evasive construct». *New Ideas in Psychology, On defining and interpreting constructs: Ontological and epistemological constraints*, 31 (1): 4-12. <https://doi.org/10.1016/j.newideapsych.2011.02.003>.
- Tabachnick, B. G., & Fidell, L. S. (2007). *Using multivariate statistics* (5th ed.). Boston, MA: Allyn & Bacon.
- Tabachnick, B. G., & Fidell, L. S. (2019). *Using multivariate statistics* (7ª ed.). Pearson.
- Thakkar, J. J. (2020). *Structural Equation Modelling: Application for Research and Practice* (with AMOS and R). Springer Singapore.

- Teo, T., Tsai, L. T., & Yang, C.-C. (2013). Applying structural equation modeling (SEM) in educational research. En M. S. Khine (Ed.), *Application of structural equation modeling in educational research and practice* (pp. 3–21). SensePublishers. [https://doi.org/10.1007/978-94-6209-332-4\\_1](https://doi.org/10.1007/978-94-6209-332-4_1).
- Teo, T., & Tsai, L. T. (Eds.). (2014). *Application of structural equation modeling in educational research and practice*. Sense Publishers. <https://doi.org/10.1007/978-94-6209-397-0>.
- Tu, Z., Yuan, Y., 2012. Understanding user's behaviors in coping with security threat of mobile devices loss and theft. *Proceedings of the 45th Hawaii International Conference on System Science*, pp. 1393–1402.
- Vance, Anthony, Mikko Siponen, y Seppo Pahnla. (2012). «Motivating IS security compliance: Insights from Habit and Protection Motivation Theory». *Information & Management* 49 (3): 190-98. <https://doi.org/10.1016/j.im.2012.04.002>.
- Waldrop, M. Mitchell. (2016). «How to Hack the Hackers: The Human Side of Cybercrime». *Nature* 533 (7602): 164-67. <https://doi.org/10.1038/533164a>.
- Watkins, M. (2021). *A step-by-step guide to exploratory factor analysis with SPSS* (1st ed.). Routledge.
- Westland, J. C. (2019). *Structural equation models: From paths to networks* (2.<sup>a</sup> ed.). Springer.
- Witte, K., Cameron, K. A., McKeon, J. K., & Berkowitz, J. M. (1996). Predicting risk behaviors: Development and validation of a diagnostic scale. *Journal of Health Communication*, 1(4), 317–341. <https://doi.org/10.1080/108107396127988>.
- Workman, Michael, William H. Bommer, y Detmar Straub. (2008). «Security lapses and the omission of information security measures: A threat control model and empirical test». *Computers in Human Behavior*, Including the Special Issue: Electronic Games and Personalized eLearning Processes, 24 (6): 2799-2816. <https://doi.org/10.1016/j.chb.2008.04.005>.
- Yousafzai, Shumaila Y., Gordon R. Foxall, y John G. Pallister. (2010). «Explaining Internet Banking Behavior: Theory of Reasoned Action, Theory of Planned Behavior, or Technology Acceptance Model?» *Journal of Applied Social Psychology* 40 (5): 1172-1202. <https://doi.org/10.1111/j.1559-1816.2010.00615.x>.

## ANEXO 1: Modelo SEM

En este estudio se empleó el Modelado de Ecuaciones Estructurales (SEM, por sus siglas en inglés), un conjunto de técnicas estadísticas multivariadas diseñadas para modelar relaciones complejas entre variables observadas y latentes (Hoyle, 2012). El SEM permite analizar simultáneamente múltiples relaciones entre variables, integrando tanto variables observadas (también llamadas manifiestas) como variables latentes, es decir, variables que no son directamente medibles, pero que se infieren a partir de múltiples indicadores observables (Pett et al., 2003).

El SEM modela estas relaciones mediante ecuaciones estructurales, que representan conexiones hipotéticas, frecuentemente de carácter causal o secuencial, entre las variables. Su uso se orienta a la evaluación de modelos teóricos: el investigador plantea un modelo basado en hipótesis derivadas de la teoría, y SEM permite verificar empíricamente si los datos respaldan dicho modelo (Ramlall, 2016). Por esta razón, el SEM se considera una técnica confirmatoria, más que exploratoria; es decir, su propósito principal es confirmar o refutar hipótesis preexistentes, en lugar de identificar relaciones emergentes sin guía teórica.

El enfoque SEM combina dos componentes fundamentales: a) El modelo de medición, que especifica qué variables observadas (ítems) miden cada factor latente y los patrones de covarianza entre estos ítems; b) El modelo estructural, que especifica las relaciones causales o correlacionales entre las variables latentes, diferenciando entre variables independientes y dependientes (Ramlall, 2016).

A continuación, se describen brevemente los distintos aspectos del Modelo de Medición y del Modelo Estructural. En el caso específico de este estudio, el Modelo de Medición incluyó el Análisis Factorial Exploratorio, Análisis de consistencia interna, Análisis Factorial Confirmatorio y Análisis de validez convergente y discriminante.

### Modelo de Medición

---

#### Análisis Factorial Exploratorio (AFE)

El Análisis Factorial Exploratorio (AFE) es una técnica estadística empleada principalmente en las etapas iniciales del desarrollo de escalas, con el objetivo de identificar y describir la estructura subyacente de un conjunto de variables observadas. A diferencia de los enfoques confirmatorios, el AFE no parte de un modelo teórico predefinido ni impone restricciones previas sobre las relaciones entre variables observadas y factores latentes. Por el contrario, permite que los datos revelen los patrones de asociación existentes (Brown, 2015).

El objetivo principal del AFE es descubrir dimensiones latentes —constructos no observables— que expliquen las covariaciones entre las variables observadas. La técnica identifica grupos de variables que están fuertemente correlacionadas entre sí, bajo el



supuesto de que comparten una causa común: un mismo factor latente. Simultáneamente, el AFE permite reducir la dimensionalidad de los datos, al sintetizar la información contenida en múltiples variables observadas en un número menor de factores. Esto facilita la interpretación de los datos y contribuye a la depuración de escalas mediante la eliminación de ítems con cargas factoriales bajas<sup>4</sup> o cargas cruzadas<sup>5</sup>. La decisión sobre qué ítems conservar se basa en criterios estadísticos y en la interpretabilidad teórica de los factores obtenidos (Fabrigar & Wegener, 2011).

### **Matriz de Cargas Factoriales (matrix of factor loadings)**

Dentro del AFE, uno de los productos centrales es la **matriz de cargas factoriales**, la cual permite identificar qué variables observadas están más estrechamente asociadas a cada factor. La matriz de cargas factoriales constituye un resultado clave del Análisis Factorial Exploratorio (AFE), ya que cuantifica la magnitud de la relación entre las variables observadas y los factores latentes extraídos. Cada carga factorial refleja el grado de asociación entre un ítem y un factor (Tabachnick & Fidell, 2007).

Cargas elevadas (por ejemplo,  $\geq 0,60$ ) indican que la variable observada está fuertemente asociada a un factor específico, lo que sugiere que representa adecuadamente el constructo latente correspondiente. Por el contrario, cargas bajas (por ejemplo,  $< 0,30$ ) reflejan una asociación débil, lo que puede implicar que el ítem no contribuye significativamente a ningún factor. Sin embargo, Aunque no existe un umbral universalmente aceptado, comúnmente se consideran significativas las cargas iguales o superiores a 0,40, aunque algunos autores proponen valores desde 0,32 como mínimo interpretativo. Además, cuando una variable presenta cargas relevantes en más de un factor (por ejemplo, superiores a 0,40 en dos o más factores), se identifican cargas cruzadas, lo cual complica la interpretación factorial y puede sugerir solapamiento conceptual entre factores (Tabachnick & Fidell, 2007).

La matriz de cargas factoriales permite, por tanto, evaluar la coherencia en el agrupamiento de ítems, detectar ítems problemáticos y tomar decisiones informadas respecto a la retención o eliminación de variables. En una solución factorial ideal, cada ítem presenta una carga elevada en un único factor y cargas mínimas en los restantes, lo que favorece la claridad interpretativa y respalda la validez estructural del instrumento (Tabachnick & Fidell, 2007).

---

<sup>4</sup> Las cargas factoriales son los coeficientes que indican la fuerza de la relación entre una variable observada y una dimensión latente. Una carga factorial baja indica que la variable no está fuertemente relacionada con la dimensión latente.

<sup>5</sup> Cargas cruzadas ocurren cuando una variable tiene cargas factoriales significativas en más de una dimensión latente, lo que significa que no miden claramente una sola dimensión latente.



## Análisis Factorial Confirmatorio (AFC)

A diferencia del Análisis Factorial Exploratorio (AFE), el Análisis Factorial Confirmatorio (AFC) no tiene como objetivo descubrir una estructura factorial, sino evaluar y confirmar una estructura de medida predefinida de constructos latentes, basada en especificaciones teóricas y evidencia empírica previa. El AFC es una técnica específica dentro del Modelado de Ecuaciones Estructurales (SEM), orientada a la evaluación de modelos de medición, que permite analizar el grado de ajuste del modelo a los datos empíricos y realizar modificaciones fundamentadas en la teoría cuando sea necesario (Brown, 2015).

En el AFC, el investigador propone un modelo a priori que especifica el número de factores latentes y el patrón de relaciones entre estos y los indicadores observados, es decir, qué ítems cargan en qué factores. Esta especificación debe estar guiada por la teoría y estudios previos (Brown, 2015).

Aunque el AFC se enfoca en las relaciones entre variables observadas y constructos latentes, los modelos SEM de mayor alcance pueden incluir además relaciones estructurales o causales entre las variables latentes, ampliando así el análisis más allá del modelo de medición (Harrington, 2008).

La evaluación del modelo en el AFC se basa en la matriz de correlaciones observada, comparándola con la matriz implícita estimada por el modelo teórico. Para ello, se emplean diversos índices de ajuste global, entre los cuales se destacan el *Comparative Fit Index* (CFI), el *Tucker-Lewis Index* (TLI) y el *Root Mean Square Error of Approximation* (RMSEA). Estos indicadores permiten evaluar en qué medida el modelo especificado reproduce las relaciones entre las variables observadas presentes en los datos (Brown, 2015; Hoyle, 2012). A la vez, las cargas factoriales estandarizadas son un resultado directo del Análisis Factorial Confirmatorio, ya que proporciona información esencial sobre la fuerza de la relación entre los indicadores observados y los constructos latentes hipotetizados (Brown, 2015).

A continuación, se explican cada uno de estos índices:

### **CFI (Comparative Fit Index)**

El *Comparative Fit Index* (CFI) es un índice ampliamente utilizado para evaluar el ajuste de modelos en el Análisis Factorial Confirmatorio (AFC) y el Modelado de Ecuaciones Estructurales (SEM). Este índice compara el modelo especificado por el investigador con un "modelo nulo", el cual asume independencia total entre las variables observadas. El CFI proporciona así una medida de ajuste relativo, y es menos sensible al tamaño de la muestra que otros índices como el chi-cuadrado (Teo et al., 2013).

Los valores del CFI oscilan entre 0 (ajuste nulo) y 1 (ajuste perfecto), siendo deseable obtener valores cercanos a 1. Por convención, se considera que un valor de  $CFI \geq 0.90$  indica un ajuste aceptable del modelo (indica que el 90% de la covariación en los datos puede ser reproducida por el modelo dado). No obstante, algunos autores sugieren un

umbral más estricto de  $CFI \geq 0.95$  como indicador de buen ajuste (Westland, 2019; Hu & Bentler, 1999). Valores elevados de CFI indican que el modelo propuesto ofrece una mejora sustancial en el ajuste en comparación con el modelo nulo.

### **TLI (Tucker-Lewis Index)**

El Índice de Tucker-Lewis (TLI), se utiliza para evaluar el ajuste de un modelo propuesto en comparación con un modelo basal, también llamado modelo nulo o de independencia, que asume que todas las variables observadas son totalmente independientes entre sí. Asimismo, puede emplearse para comparar modelos alternativos. Al igual que el Índice de Ajuste Comparativo (CFI), el TLI pertenece a la categoría de índices de ajuste incremental, los cuales valoran la mejora en el ajuste del modelo propuesto respecto al modelo de referencia (Hoyle, 2012; Kline, 2016).

Aunque su interpretación es similar a la del CFI, el TLI se considera un índice no normado, lo que significa que sus valores pueden exceder el rango entre 0 y 1. Sin embargo, valores cercanos a 1.0 suelen interpretarse como indicativos de un buen ajuste del modelo (Brown, 2015; Kaplan, 2008). Diversos autores sugieren que un valor de  $TLI \geq 0.95$  representa un ajuste adecuado entre los datos observados y la estructura teórica especificada (Teo et al., 2013).

El TLI, al igual que el Error Cuadrático Medio de Aproximación (RMSEA), incorpora una penalización por complejidad del modelo, penalizando la inclusión de parámetros libremente estimados que no mejoran significativamente el ajuste. Por ello, se considera un índice que ajusta el modelo en función de la parsimonia (Brown, 2015; Kline, 2016). Finalmente, estudios de simulación (Monte Carlo) han demostrado que el TLI es relativamente poco sensible al tamaño muestral, lo que lo convierte en una medida robusta en contextos con tamaños de muestra moderados o reducidos (Maruyama, 1998).

### **RMSEA (Root Mean Square Error of Approximation)**

El Error Cuadrático Medio de Aproximación (RMSEA) estima el error de aproximación poblacional, permitiendo cierto grado de discrepancia entre el modelo hipotetizado y la matriz de covarianza poblacional. Esto lo convierte en una herramienta útil para evaluar si el modelo proporciona un ajuste razonablemente bueno en lugar de perfecto (Brown, 2015; O'Rourke & Hatcher, 2013).

El RMSEA es un índice de mal ajuste (*badness-of-fit*), lo que significa que valores más bajos indican un mejor ajuste. Un valor de 0,00 representa ajuste perfecto (Beaujean, 2014). A la vez, el RMSEA penaliza la complejidad del modelo, al estimar la discrepancia por grado de libertad, favoreciendo así modelos más parsimoniosos cuando el ajuste es similar (Beaujean, 2014; Fabrigar & Wegener, 2012).

La interpretación del RMSEA, según Browne & Cudeck (1993) y Hu & Bentler (1999):  $\leq 0.05$ : Ajuste cercano o excelente; 0.05 a 0.08: Ajuste aceptable o razonable; 0.08 a 0.10: Ajuste mediocre o marginal;  $> 0.10$ : Ajuste pobre o inadecuado.

### **Cargas Factoriales Estandarizadas**

Las cargas factoriales estandarizadas<sup>6</sup> proporcionan información directa sobre la fuerza de la relación entre un constructo latente y sus indicadores observados, y son cruciales para determinar si existe validez convergente dentro de un análisis confirmatorio (Mehmetoglu & Venturini, 2021).

Las cargas factoriales estandarizadas representan la correlación entre el constructo latente y su indicador correspondiente. Puede interpretarse como el cambio esperado en unidades de desviación estándar en la variable observada por cada cambio de una unidad de desviación estándar en el factor latente, asumiendo que todas las demás variables se mantienen constantes (Brown, 2015).

Las cargas estandarizadas ayudan a evaluar la importancia relativa de cada variable observada para la definición de su factor latente. Cargas altas (en valor absoluto, generalmente  $\geq 0.70$ , aunque este umbral puede variar según el contexto y la disciplina) indican que el indicador comparte una proporción sustancial de varianza con el constructo latente. Esto es un indicio fuerte de validez convergente, ya que sugiere que varios indicadores del mismo constructo están altamente relacionados con él (Mehmetoglu & Venturini, 2021).

Es crucial considerar que la significancia estadística de la carga factorial (generalmente evaluada mediante un *p-value*) es tan importante como la magnitud de la carga. Una carga factorial puede ser sustancial, pero si no es estadísticamente significativa, su fiabilidad puede ser cuestionable. Es decir, La significación estadística proporciona evidencia de que la relación entre el indicador y el factor es probablemente real y no producto del azar (Brown, 2015).

Finalmente, es importante destacar que el análisis de cargas factoriales se aplica tanto en análisis factorial exploratorio (AFE) como confirmatorio (AFC), con diferencias clave en su enfoque. En el AFE, el investigador no especifica el número de factores ni las relaciones entre variables y factores, permitiendo que los datos revelen la estructura subyacente. En contraste, el AFC evalúa la significancia de las cargas factoriales dentro de un modelo teórico predefinido, determinando si las relaciones hipotetizadas son consistentes con los datos (Brown, 2015; Watkins, 2021).

---

<sup>6</sup> Las cargas factoriales no estandarizadas están en la escala original de las variables observadas. Esto puede dificultar la comparación de la importancia relativa de diferentes indicadores para un mismo factor, especialmente si las variables observadas tienen escalas de medición diferentes (Brown, 2015).



## Consistencia interna, Análisis de Validez Convergente y Análisis de Validez Discriminante

La evaluación de la validez de constructo se fundamenta en tres procesos complementarios: la consistencia interna, la validez convergente y la validez discriminante. La consistencia interna busca determinar si los ítems que componen un mismo factor presentan una covariación significativa entre sí, lo que indicaría que están midiendo de manera coherente un mismo constructo subyacente. En tanto, la validez convergente se manifiesta en la correlación elevada entre mediciones que apuntan al mismo constructo teórico, demostrando así su congruencia interna. Por otra parte, la validez discriminante se enfoca en asegurar la independencia entre constructos distintos, evidenciando la ausencia de correlaciones significativas entre sus respectivas mediciones. Estos análisis son cruciales para establecer la solidez del constructo medido (Brown, 2015).

### Análisis de Consistencia Interna

Una vez identificados los factores y sus indicadores correspondientes, es fundamental evaluar la **consistencia interna** de cada dimensión latente. Este análisis busca determinar si los ítems que componen un mismo factor presentan una covariación significativa entre sí, lo que indicaría que están midiendo de manera coherente un mismo constructo subyacente (Hoyle, 2023; Mehmetoglu & Venturini, 2021). Para evaluar esta consistencia, se utilizan indicadores como el **alfa de Cronbach**, el análisis de **correlaciones inter-ítem** y el indicador de confiabilidad compuesta (Kline, 2016; Brown, 2015). Una alta consistencia interna refuerza la fiabilidad de la medición y la validez de los factores obtenidos.

Es importante destacar que, en el contexto de modelos complejos como el Modelado de Ecuaciones Estructurales (SEM), una adecuada consistencia interna es esencial. Si los indicadores que conforman un constructo no muestran un nivel aceptable de consistencia, ello puede comprometer la validez de las inferencias realizadas sobre las relaciones estructurales del modelo (Hoyle, 2023).

### Alfa de Cronbach (Cronbach's Alpha)

El alfa de Cronbach ( $\alpha$ ) es una estimación de la consistencia interna, una forma de fiabilidad que evalúa el grado en que los ítems de una escala presentan respuestas consistentes entre sí. Una consistencia interna baja puede indicar que los ítems son demasiado heterogéneos en contenido, lo que sugeriría que no todos están midiendo el mismo constructo latente. Por otro lado, Un valor alto de alfa sugiere que la escala presenta una buena consistencia interna, lo cual respalda la interpretación de que los ítems están midiendo aspectos relacionados de un mismo constructo. No obstante, valores excesivamente altos (por ejemplo,  $> 0.95$ ) podrían señalar redundancia entre ítems. Por el contrario, un valor bajo de alfa puede indicar que algunos ítems no están contribuyendo de manera coherente a la medición del constructo (Brown, 2015; Kline, 2016).



Los valores del alfa de Cronbach suelen interpretarse según los siguientes rangos (George & Mallery, 2003; DeVellis, 2016):  $\alpha \geq 0.90$ : Excelente confiabilidad;  $0.80 \leq \alpha < 0.90$ : Buena confiabilidad;  $0.70 \leq \alpha < 0.80$ : Aceptable;  $0.60 \leq \alpha < 0.70$ : Dudosa;  $\alpha < 0.60$ : No confiable.

Es importante considerar que el alfa de Cronbach asume la unidimensionalidad de la escala (esto significa que se espera que el conjunto de ítems mida un único constructo latente); por lo tanto, se recomienda realizar análisis factoriales exploratorios o confirmatorios previos para verificar esta suposición. Asimismo, el coeficiente puede verse influido por el número de ítems, lo que puede inflar artificialmente su valor en escalas muy extensas (DeVellis, 2016; Netemeyer et al., 2003).

### **Correlación inter-ítems**

La correlación inter-ítems es una herramienta fundamental para comprender las relaciones entre los ítems de una escala, así como para contribuir a la evaluación de su consistencia interna y la calidad de la medición (Hoyle, 2012). Esta medida refleja el grado de asociación entre las respuestas a diferentes ítems y proporciona información clave sobre la homogeneidad del contenido de la escala (DeVellis, 2016).

Cuando los ítems están diseñados para evaluar un atributo subyacente común, se espera que las correlaciones entre ellos sean sustanciales, indicando que están midiendo el mismo constructo (Hoyle, 2012; DeVellis, 2016). Por el contrario, ítems con correlaciones bajas respecto a los demás pueden no estar alineados con el constructo teórico y, por tanto, podrían ser candidatos para su revisión o eliminación (DeVellis, 2016).

No existe un umbral único e inamovible para considerar adecuada una correlación inter-ítems, ya que esta dependerá del tipo de constructo evaluado y del contexto del estudio (Netemeyer et al., 2003). Por ejemplo, para constructos amplios o abstractos, es razonable esperar correlaciones más modestas, mientras que constructos más estrechamente definidos deberían reflejar asociaciones más fuertes entre ítems (DeVellis, 2016). Sin embargo, se reconocen ciertas directrices comunes (Pett et al., 2003): i) Las correlaciones deben ser positivas: correlaciones negativas pueden indicar errores en la redacción, como ítems invertidos que no han sido adecuadamente recodificados, o incluso la presencia de constructos distintos; ii) Correlaciones bajas pueden reflejar falta de cohesión conceptual entre ítems y sugieren que no están midiendo adecuadamente el mismo constructo; iii) Correlaciones muy altas (por ejemplo, superiores a 0,80) pueden indicar redundancia entre ítems, es decir, contenido solapado que no aporta valor adicional a la medición.

En general, una correlación inter-ítems promedio inferior a 0.15 puede ser motivo de preocupación, mientras que un rango entre 0.20 y 0.50 suele ser aceptable, dependiendo de la amplitud conceptual del constructo (Clark y Watson, 1995; Robinson et al., 1991; Everitt y Skrandal, 2010).

Finalmente, Clark y Watson (1995) sugieren que la correlación inter-ítems promedio puede ser un indicador más útil de consistencia interna que el coeficiente alfa de Cronbach, ya que este último puede verse afectado por el número de ítems y otras limitaciones estadísticas.

### **Confiabilidad Compuesta (CR)**

La Confiabilidad Compuesta (CR) es una medida que nos dice qué tan bien un conjunto de indicadores mide un mismo concepto subyacente (constructo latente). En otras palabras, nos dice si las respuestas a esas preguntas son consistentes y confiables (Brown, 2015).

El indicador se calcula usando las cargas factoriales. La fórmula compara la varianza explicada por el constructo con la varianza de error de los ítems. Indica cuanta varianza de los ítems se explica por el constructo latente subyacente, en relación con la varianza del error. La fórmula para CR considera la varianza al cuadrado de la suma de las cargas factoriales dividida por la varianza al cuadrado de la suma de las cargas factoriales más la suma de las varianzas del error de los ítems (Brown, 2015).

Una CR alta sugiere que los indicadores de un constructo miden consistentemente el mismo concepto subyacente (las preguntas miden de forma consistente el mismo concepto), lo cual es un requisito previo para la validez convergente. Sin embargo, una CR alta es necesaria, pero no suficiente. También necesitamos evaluar la "Varianza Media Extraída (AVE)", que nos dice qué tanta varianza del concepto es explicada por las preguntas (O'Rourke & Hatcher, 2013; Brown, 2015).

### **Validez Convergente**

La validez convergente se refiere al grado en que múltiples indicadores de un mismo constructo están altamente correlacionados entre sí, lo que sugiere que efectivamente miden el mismo concepto subyacente. En otras palabras, evalúa si los ítems que deberían estar relacionados lo están de manera significativa (Brown, 2006).

Su evaluación suele basarse en los siguientes indicadores: Varianza media extraída (AVE) y Confiabilidad Compuesta (CR).

### **Varianza Media Extraída (AVE)**

La Varianza Media Extraída (AVE, *Average Variance Extracted*) es una medida estadística que se usa para evaluar si un conjunto de ítems (indicadores) realmente está midiendo el mismo concepto o idea teórica, es decir, si tienen validez convergente.

La AVE calcula cuánta parte de la información (varianza) de los indicadores es explicada por el constructo al que pertenecen. Se obtiene promediando los cuadrados de las cargas factoriales (es decir, cuánto pesa cada ítem en el constructo). En otras palabras, se calcula sumando el cuadrado de cada carga factorial de los ítems que miden un constructo y dividiendo esa suma por el número de ítems (Hair et al., 2017).



Un valor de AVE de 0.50 o superior se considera generalmente aceptable y proporciona evidencia de una adecuada validez convergente a nivel del constructo. Si la AVE es 0.50 o más, significa que el constructo explica al menos el 50% de la varianza de sus indicadores (Teo & Tsai, 2014).

Un AVE inferior a 0.50 indica que, en promedio, queda más varianza en el error de los ítems que en la varianza explicada por el constructo. Si la AVE es menor a 0.50, quiere decir que hay más ruido o error que información útil en los indicadores, lo que sugiere baja validez convergente (Teo & Tsai, 2014).

## **Validez Discriminante**

La validez discriminante se refiere a la capacidad de un constructo para diferenciarse claramente de otros constructos en un modelo. Es decir, garantiza que cada constructo mida un concepto único y distinto, no coincidiendo sustancialmente con otros constructos teóricamente diferentes. Por lo tanto, dos constructos diferentes no deberían presentar correlaciones extremadamente altas entre sí (por ejemplo, correlaciones estandarizadas iguales o superiores a 0,85), ya que esto indicaría un problema de validez discriminante (Brown, 2015).

Existen dos criterios ampliamente utilizados para la evaluación de la validez discriminante:

### **Criterio de Fornell-Larcker**

Este criterio establece que la raíz cuadrada del Average Variance Extracted (AVE) para cada constructo latente debe ser mayor que las correlaciones entre dicho constructo y todos los demás constructos del modelo. La lógica es sencilla: cada constructo debería explicar mejor la varianza de sus propios indicadores (ítems o preguntas) que la varianza que comparte con otros constructos, lo cual respalda su distinción empírica. Así, al construir una matriz que incluya las raíces cuadradas del AVE en la diagonal y las correlaciones entre constructos fuera de la diagonal, la validez discriminante se confirma si cada valor diagonal es superior a los valores correspondientes en sus filas y columnas (Hair et al., 2017).

### **Criterio de cargas cruzadas**

Este criterio plantea que la carga factorial de un indicador en su propio constructo debe ser mayor que cualquier carga que dicho indicador tenga en otros constructos. Esto indica que cada indicador está más asociado al constructo específico que pretende medir, en lugar de medir indirectamente otros constructos (Hair et al., 2017).



## Modelo Estructural

---

El modelo estructural es el núcleo del modelado de ecuaciones estructurales (SEM), y su función principal es representar las relaciones hipotetizadas entre variables, tanto latentes como observadas. Su propósito es probar empíricamente teorías sobre la dependencia o causalidad entre constructos. Está compuesto por variables y caminos que cuantifican la magnitud de sus relaciones mediante parámetros estimables (Thakkar, 2020).

Tiene las siguientes etapas de desarrollo:

- **Especificación del modelo estructural**

La **especificación del modelo** consiste en definir claramente los constructos, es decir, las variables latentes que representan conceptos teóricos, formular hipótesis sobre las relaciones causales esperadas entre estos constructos y presentar visualmente estas relaciones mediante un diagrama de trayectorias (path diagram), que muestra gráficamente cómo interactúan las variables latentes y observadas dentro del modelo propuesto (Hoyle, 2012).

El modelo estructural se representa comúnmente mediante diagramas de caminos y ecuaciones, e incluye relaciones directas (una variable influye directamente en otra), indirectas (el efecto de una variable sobre otra se da a través de una tercera variable, denominada mediadora o interviniente) y moderadas (es una tercera variable que influye en la naturaleza de la relación entre una variable independiente y una variable dependiente, modificando la fuerza o la dirección de ese efecto) (Thakkar, 2020; Hair et al., 2017).

Las relaciones entre estas variables se representan mediante caminos en el diagrama de ruta. Estos caminos indican la dirección hipotetizada de la influencia (con flechas) o la simple covariación (con flechas de doble punta, aunque estas son más comunes en la parte de medición o entre variables exógenas en el modelo estructural) (Hoyle, 2012). Cada camino en el modelo estructural está asociado con un parámetro, que cuantifica la magnitud y dirección de la relación hipotetizada. Estos parámetros son típicamente coeficientes de regresión (también llamados coeficientes de ruta o coeficientes estructurales) (Thakkar, 2020).

Estos parámetros son desconocidos a priori y se estiman a partir de los datos de la muestra utilizando algoritmos de optimización para minimizar la discrepancia entre la matriz de covarianzas observada y la matriz de covarianzas implicada por el modelo (Raykov & Marcoulides, 2006).

Junto con las estimaciones de los parámetros, se obtienen errores estándar que permiten evaluar la significancia estadística de cada relación hipotetizada. Es en esta etapa donde la consideración de posibles falsos positivos (error de Tipo I) es crucial, especialmente si se están probando múltiples relaciones simultáneamente (Thakkar, 2020).

- **Identificación del modelo**

La identificación del modelo implica verificar que exista suficiente información estadística (grados de libertad adecuados) para estimar sus parámetros de manera única. Si en esta etapa el modelo no resulta identificado, es necesario realizar ajustes en su estructura para asegurar que pueda ser estimado correctamente y producir soluciones válidas (Kline, 2016).

Para que el modelo sea estimable, debe estar identificado, lo que implica que haya suficientes datos para calcular los parámetros de manera única. Un modelo es estimable cuando es posible obtener valores numéricos para cada uno de sus parámetros libres utilizando los datos disponibles. Los parámetros libres son aquellos que el investigador desea estimar a partir de los datos, como las cargas factoriales, las varianzas, las covarianzas y los coeficientes de regresión (Kline, 2016; Schumacker & Lomax, 2016).

En términos prácticos, para que exista suficiente información para calcular estos parámetros de forma única, la cantidad de elementos conocidos en la matriz empírica de varianzas y covarianzas (cuya cantidad se determina por la fórmula  $p(p+1)/2$ , siendo  $p$  el número de variables observadas) debe ser igual o mayor al número de parámetros desconocidos que se pretenden estimar. Dependiendo de esta relación, un modelo puede clasificarse en tres estados fundamentales de identificación: subidentificado, justo identificado o sobreidentificado (Schumacker & Lomax, 2016).

Además, la regla de conteo, que establece que el número de datos debe ser mayor o igual que el número de parámetros, es una condición necesaria pero no suficiente para garantizar la identificación completa del modelo. Existen condiciones adicionales fundamentales que deben cumplirse para asegurar la identificación, tales como: el escalamiento adecuado de variables latentes (ya sea fijando la varianza del factor o una carga factorial de referencia), una estructura clara del modelo de medición con suficientes indicadores por factor (idealmente tres o más para evitar problemas empíricos), y una estructura adecuada del modelo estructural que evite dependencias lineales o relaciones recíprocas no identificables. También debe evitarse la subidentificación empírica derivada de patrones específicos en los datos (como matrices de covarianzas demasiado débiles o relaciones estructurales insuficientes). (Brown, 2015; Kline, 2016; Thakkar, 2020).

En resumen, para que un modelo de ecuaciones estructurales (SEM) sea estimable y produzca parámetros interpretables de manera significativa y confiable, debe estar correctamente identificado. Esto implica asegurar suficiente información en los datos observados, en relación con la complejidad teórica del modelo, para obtener soluciones únicas y válidas para cada parámetro libre. Un modelo no identificado conlleva el riesgo significativo de producir resultados arbitrarios e inválidos desde un punto de vista científico.

- **Estimación de Parámetros**

Se aplican métodos estadísticos (como Máxima Verosimilitud o PLS) para estimar los coeficientes de trayectoria, que indican magnitud y dirección de las relaciones entre variables latentes.

En el contexto del Modelado de Ecuaciones Estructurales (SEM), la estimación de parámetros consiste en calcular los valores numéricos que representan las relaciones entre variables latentes, con base en los datos observados. Este proceso busca minimizar la discrepancia entre la matriz de covarianzas observada y la que implica el modelo teórico. Para ello, se emplean métodos estadísticos como la Máxima Verosimilitud (ML), que asume normalidad multivariada y estima todos los parámetros simultáneamente, o PLS-SEM (Mínimos Cuadrados Parciales), orientado a la predicción y más flexible frente a supuestos distribucionales. El resultado de esta estimación son los coeficientes de trayectoria, que indican tanto la magnitud (fuerza) como la dirección (positiva o negativa) de las relaciones entre variables latentes. Estos coeficientes permiten evaluar empíricamente si las relaciones propuestas en el modelo estructural se sostienen, considerando también la calidad del modelo de medición. La elección del método de estimación dependerá del objetivo del estudio (ajuste teórico o predicción) y de las características de los datos disponibles (Teo & Tsai, 2014, Hair et al., 2017).

- **Evaluación del ajuste del modelo estructural**

La evaluación del ajuste del modelo implica utilizar índices estadísticos específicos, como RMSEA, CFI y TLI, para determinar en qué medida el modelo propuesto representa adecuadamente los datos observados y verificar si la estructura planteada es consistente con la evidencia empírica disponible.

- **Interpretación de resultados**

La interpretación de resultados implica analizar la significancia estadística y la dirección de los coeficientes de trayectoria para determinar si las relaciones planteadas entre las variables del modelo se sostienen empíricamente y si estas relaciones son coherentes con las hipótesis teóricas formuladas previamente. La significancia se evalúa mediante valores  $p$ , estadísticos  $t$  e intervalos de confianza, y permite determinar si dichas relaciones son estadísticamente válidas. La dirección, indicada por el signo del coeficiente, debe coincidir con lo esperado teóricamente (positiva o negativa). Este análisis conjunto permite verificar si las relaciones propuestas se sostienen empíricamente y si son coherentes con las hipótesis formuladas. Incluso cuando el modelo muestra un buen ajuste global, es esencial revisar cada relación específica, ya que algunas pueden no ser significativas o presentar direcciones opuestas a las esperadas, lo cual puede sugerir la necesidad de ajustar la teoría, considerar variables omitidas o repensar el modelo. Por tanto, la interpretación de resultados es un proceso crítico que exige tanto rigor estadístico como comprensión teórica del fenómeno estudiado (Teo & Tsai, 2014, Brown, 2006).

## ANEXO 2: Variables Latentes

A continuación, se presentan los diferentes ítems utilizados inicialmente para la definición de cada constructo.

### Teoría de Motivación de Protección

- **Evaluación de amenazas a la seguridad:** La evaluación de un individuo sobre el nivel de peligro que representa un evento de seguridad (Crossler, 2010).

**Vulnerabilidad de seguridad percibida:** La vulnerabilidad percibida es la evaluación que hace un usuario sobre la probabilidad de que ocurra un evento de seguridad amenazante (Crossler, 2010). En este caso, se refiere a la medida en que un cliente cree que será víctima de fraude en la banca en línea.

| N | Item                                                                                      |
|---|-------------------------------------------------------------------------------------------|
| 1 | Estoy en riesgo de ser víctima de fraude informático                                      |
| 2 | Es posible que me convierta en víctima de fraude informático                              |
| 3 | Sé que podría ser vulnerable a violaciones de seguridad si no sigo las recomendaciones    |
| 4 | Mi datos y sitios financieros son vulnerables a violaciones de seguridad                  |
| 5 | Mis sitios financieros pueden verse comprometidos si no sigo recomendaciones de seguridad |
| 6 | Creo que podría ser vulnerable al fraude informático                                      |

Nota: Items adaptados de Jansen y Van Schaik (2018b), Ifinedo (2012), Lee (2011), Witte (1996).

**Gravedad percibida:** es la evaluación subjetiva de un usuario sobre la magnitud de las consecuencias negativas derivadas de un evento de seguridad amenazante. En el contexto de la banca en línea, este estudio la aplica a la percepción de un cliente sobre la seriedad de las consecuencias del fraude (Crossler, 2010).

| N | Item                                                                                                          |
|---|---------------------------------------------------------------------------------------------------------------|
| 1 | Que alguien acceda al sitio web de mi banco sin mi consentimiento o conocimiento es un problema grave para mí |
| 2 | La pérdida financiera resultante de un fraude informático es un problema grave para mí                        |
| 3 | Creo que el fraude informático es un problema grave                                                           |
| 4 | El fraude informático podría dañar gravemente mi situación financiera                                         |
| 5 | Es muy probable que un fraude informático me cause problemas importantes                                      |
| 6 | Que alguien ataque mis sitios financieros sería muy perjudicial para mí                                       |

Nota: Items adaptados de Ifinedo (2012), Lee (2011), Witte (1996).

- **Evaluación del afrontamiento de seguridad:** La evaluación de una persona sobre su capacidad para realizar una acción específica y su confianza en que esa acción evitará pérdidas o daños ante una amenaza de seguridad, considerando que el costo percibido no sea excesivo (Crossler, 2010).

**Autoeficacia:** La confianza de un individuo en su capacidad para realizar el comportamiento recomendado para prevenir o mitigar el evento de seguridad amenazante (Crossler, 2010).

| N | Item                                                                                                    |
|---|---------------------------------------------------------------------------------------------------------|
| 1 | Cuento con las habilidades necesarias para seguir las recomendaciones de seguridad                      |
| 2 | Cuento con los conocimientos necesarios para seguir las recomendaciones de seguridad                    |
| 3 | Podría ser capaz de seguir las recomendaciones de seguridad incluso sin nadie cerca que me muestre cómo |
| 4 | Tengo la experiencia para implementar medidas preventivas que eviten los fraudes informáticos           |
| 5 | Soy capaz de cumplir con las recomendaciones de seguridad                                               |
| 6 | Podría adoptar las recomendaciones de seguridad sin ayuda de terceros                                   |

Nota: Items adaptados de Chen et al. (2018), George (2004), Witte (1996).

**Eficacia de la respuesta:** Se define como la percepción de una persona sobre la efectividad de una acción para reducir o eliminar una amenaza. En este estudio, se refiere a la confianza en que un comportamiento recomendado prevendrá o mitigará un evento de seguridad amenazante (Crossler, 2010).

| N | Item                                                                                                                                        |
|---|---------------------------------------------------------------------------------------------------------------------------------------------|
| 1 | Las recomendaciones de seguridad ayudan a prevenir el fraude informático                                                                    |
| 2 | Cumplir las recomendaciones de seguridad es efectivo para prevenir el fraude informático                                                    |
| 3 | Si sigo las recomendaciones de seguridad, es menos probable que sea víctima de un fraude                                                    |
| 4 | Adoptar las recomendaciones de seguridad es una forma eficaz de disuadir los ataques de piratas informáticos                                |
| 5 | Seguir las recomendaciones de seguridad evitará que los piratas informáticos obtengan acceso a información personal o financiera importante |
| 6 | Las medidas preventivas disponibles para evitar que las personas obtengan acceso a mis sitios financieros son adecuadas                     |

Nota: Items adaptados de Jansen y Van Schaik (2018b), Ifinedo (2012), Witte (1996).

**Costos de la respuesta:** Se definen como la percepción del usuario sobre el esfuerzo cognitivo, tiempo, dinero u otros recursos necesarios para implementar una acción de protección (Milne et al., 2000; Menard et al., 2017). Según Crossler (2010), las medidas de precaución se adoptan cuando sus costos percibidos son menores que las posibles pérdidas por una amenaza. Es decir, este factor enfatiza los costes de oportunidad percibidos en términos de dinero, tiempo y esfuerzo invertidos en adoptar el

comportamiento recomendado. Por lo tanto, el costo percibido influye directamente en la disposición del usuario a actuar ante riesgos de seguridad de la información.

| <b>N</b> | <b>Item</b>                                                                                  |
|----------|----------------------------------------------------------------------------------------------|
| 1        | Se necesita mucho esfuerzo para protegerse del fraude informático                            |
| 2        | Adoptar las recomendaciones de seguridad requiere mucho esfuerzo mental                      |
| 3        | Cumplir las recomendaciones de seguridad requeriría comenzar un nuevo hábito                 |
| 4        | Se necesita mucho tiempo y esfuerzo para seguir las recomendaciones de seguridad             |
| 5        | No seguiría las recomendaciones de seguridad porque toma demasiado tiempo                    |
| 6        | Seguir las recomendaciones de seguridad requeriría un esfuerzo considerable además de tiempo |

Nota: Items adaptados de Jansen y van Schaik (2018b), Lee (2011), Ophoff y Lakay (2019), Vance et al. (2012), Ng et al. (2009), Ifinedo (2012).

### **Teoría del Comportamiento Planificado**

**Normas subjetivas:** Definen la percepción de un individuo sobre las creencias de las personas relevantes para él en relación con un comportamiento específico, en este caso, tomar medidas de seguridad contra el fraude informático (Ifinedo, 2012). Esencialmente, reflejan la creencia del individuo sobre la aprobación o desaprobación de otros significativos hacia su participación en dicha conducta.

| <b>N</b> | <b>Item</b>                                                                                                 |
|----------|-------------------------------------------------------------------------------------------------------------|
| 1        | Mis familiares y amigos piensan que debería tomar medidas de seguridad para protegerme del fraude           |
| 2        | Las personas importantes para mí pensarían que debo tomar medidas de seguridad contra el fraude informático |
| 3        | Mis compañeros de trabajo piensan que debería tomar medidas de seguridad para protegerme del fraude         |
| 4        | Mucha gente como yo sigue las recomendaciones de seguridad                                                  |
| 5        | Se espera de mí que siga las recomendaciones de seguridad                                                   |
| 6        | Me siento obligado por los demás a seguir las recomendaciones de seguridad contra el fraude informático     |

Nota: Items adaptados de Jansen y van Schaik (2018b). Anderson and Agarwal (2010).

**Actitud:** Se refiere al grado en que una persona evalúa favorable o desfavorablemente un comportamiento de interés, en este caso, seguir las recomendaciones de seguridad, considerando los resultados de llevarlas a cabo. En otras palabras, son los sentimientos positivos o negativos del individuo hacia la participación en esta conducta específica (Ifinedo, 2012).

| <b>N</b> | <b>Item</b>                                                     |
|----------|-----------------------------------------------------------------|
| 1        | Para mí, seguir las recomendaciones de seguridad es valioso     |
| 2        | Para mí, seguir las recomendaciones de seguridad es beneficioso |



- 3 Para mí, seguir las recomendaciones de seguridad es bueno
- 4 Para mí, seguir las recomendaciones de seguridad es útil
- 5 Para mí, seguir las recomendaciones de seguridad es algo inteligente
- 6 Para mí, seguir las recomendaciones de seguridad es lo correcto

---

Nota: Items adaptados de Ifinedo (2012).

**Control conductual percibido:** Se refiere a la percepción de una persona sobre la facilidad o dificultad de realizar el comportamiento de interés. Según Ifinedo (2012), esta percepción varía según las situaciones y acciones, lo que resulta en que una persona tenga percepciones variables del control conductual según la situación. Mientras que el énfasis de la autoeficacia está en si los individuos sienten que tienen las habilidades y capacidades adecuadas para lograr un objetivo, el control conductual percibido comprende una expresión más interactiva de la relación entre un individuo y su entorno (Workman et al., 2008).

| N | Item                                                                   |
|---|------------------------------------------------------------------------|
| 1 | Para mí seguir las recomendaciones de seguridad es fácil               |
| 2 | Seguir las recomendaciones de seguridad no es difícil                  |
| 3 | Si quisiera, podría fácilmente seguir las recomendaciones de seguridad |
| 4 | Seguir las recomendaciones de seguridad no es complicado               |
| 5 | Es fácil protegerse del fraude informático                             |
| 6 | Para mí, seguir las recomendaciones de seguridad no es algo difícil    |

---

Nota: Items adaptados de Workman et al. (2008), Tu & Yuan (2012).

