



POLITECNICO
MILANO 1863

SCUOLA DI INGEGNERIA INDUSTRIALE
E DELL'INFORMAZIONE

Understanding cybervictimization through smartphone addiction, online gambling and routine activities: an integrated approach using PLS-SEM and Machine Learning

TESI DI LAUREA MAGISTRALE IN
Management Engineering - Ingegneria Gestionale

Author: **Federico Trevisan**

Student ID: **10787720**

POLIMI Advisor: **Sergio Terzi**

UPM Advisor: **Alberto Urueña López**

Academic Year: **2025-26**

Federico Trevisan

Master Thesis - Trabajo de Fin de Máster

Understanding cybervictimization through smartphone addiction, online gambling and routine activities: an integrated approach using PLS-SEM and Machine Learning



POLITECNICO
MILANO 1863

SCUOLA DI INGEGNERIA INDUSTRIALE
E DELL'INFORMAZIONE

Scuola di Ingegneria Industriale e dell'Informazione
Politecnico di Milano, Italia

—



UNIVERSIDAD
POLITÉCNICA
DE MADRID



Departamento de Ingeniería de Organización, Administración de Empresas y
Estadística

Escuela Técnica Superior de Ingenieros Industriales
Universidad Politécnica de Madrid, España

POLIMI Advisor: **Sergio Terzi**

UPM Advisor: **Alberto Urueña López**

March 2026

Acknowledgments

Ai professori Alberto Urueña López e Sergio Terzi, per avermi guidato e consigliato durante quest'ultimo atto del mio percorso accademico.

Alla mia famiglia, per non avermi mai fatto mancare l'opportunità di crescere ed essere me stesso.

Ai miei amici, vicini e lontani, per il sostegno costante, la presenza silenziosa e i momenti di leggerezza che hanno reso questo percorso più sereno.

Abstract

The increasing diffusion of digital technologies and the intensification of online activities have led to a significant rise in cybervictimization phenomena, making it essential to understand the behavioural and situational factors influencing risk exposure. This Master Thesis analyses the role of smartphone addiction, online gambling behaviour, and digital routines in explaining vulnerability to cybervictimization, integrating Lifestyle-Routine Activity Theory (L-RAT) within a unified empirical framework.

The study is based on a national dataset and adopts an integrated methodological approach combining Partial Least Squares Structural Equation Modelling (PLS-SEM) and Machine Learning techniques. PLS-SEM is employed to test the structural relationships among theoretical constructs, modelling L-RAT as a second-order formative construct, while a Random Forest model is used to assess predictive performance and variable importance.

The results indicate that smartphone addiction and specific L-RAT dimensions, particularly Proximity to motivated offenders, are significantly associated with cybervictimization risk. Online gambling behaviour emerges as a factor that amplifies routine-based exposure mechanisms, whereas social support plays a protective role. The integration of explanatory and predictive modelling reveals convergence between theoretical relevance and predictive importance of key variables.

The study provides a dual contribution: it reinforces the applicability of Lifestyle-Routine Activity Theory in contemporary digital environments and demonstrates the value of integrating structural modelling with machine learning algorithms for the analysis of online behavioural risk.

Key-words:

Cybervictimization; Smartphone Addiction; Online Gambling; Lifestyle-Routine Activity Theory; PLS-SEM; Machine Learning; Random Forest; Digital Behaviour

UNESCO codes:

1203 Computer Science

1203.12 Databases

1209 Statistics

1209.01 Analytical Statistics

1209.03 Data Analysis

1209.13 Statistical Inference Techniques

1209.14 Statistical Prediction Techniques

6114 Social Psychology

6310 Social issues

Abstract in italiano

La crescente diffusione delle tecnologie digitali e l'intensificarsi delle attività online hanno determinato un aumento significativo dei fenomeni di cybervittimizzazione, rendendo necessario comprendere i fattori comportamentali e situazionali che ne influenzano il rischio. Questa Tesi Magistrale analizza il ruolo della dipendenza da smartphone, del comportamento di gioco d'azzardo online e delle routine digitali nella spiegazione della vulnerabilità alla cybervittimizzazione, integrando la Lifestyle-Routine Activity Theory (L-RAT) in un modello empirico unificato.

L'analisi si basa su un dataset nazionale e adotta un approccio metodologico integrato che combina Partial Least Squares Structural Equation Modeling (PLS-SEM) e Machine Learning. Il PLS-SEM è utilizzato per testare le relazioni strutturali tra i costrutti teorici, modellando la L-RAT come costrutto di secondo ordine formativo, mentre un modello Random Forest è impiegato per valutare la capacità predittiva delle variabili considerate.

I risultati mostrano che la dipendenza da smartphone e alcune dimensioni della L-RAT, in particolare la Proximity to motivated offenders, rappresentano fattori significativamente associati al rischio di cybervittimizzazione. Il gioco d'azzardo online emerge come elemento che amplifica i meccanismi di esposizione routinaria al rischio, mentre il supporto sociale assume un ruolo protettivo. L'integrazione tra modellazione esplicativa e predittiva evidenzia una convergenza tra rilevanza teorica e importanza predittiva delle variabili chiave.

Il contributo della ricerca è duplice: da un lato, rafforza l'applicabilità della L-RAT nei contesti digitali contemporanei; dall'altro, dimostra il valore dell'integrazione tra approcci strutturali e algoritmi di apprendimento automatico per l'analisi del rischio comportamentale online.

Parole chiave:

Cybervittimizzazione; Dipendenza da smartphone; Gioco d'azzardo online; Lifestyle-Routine Activity Theory; PLS-SEM; Machine Learning; Random Forest; Comportamento digitale

Codici UNESCO:

1203 Computer Science

1203.12 Databases

1209 Statistics

1209.01 Analytical Statistics

1209.03 Data Analysis

1209.13 Statistical Inference Techniques

1209.14 Statistical Prediction Techniques

6114 Social Psychology

6310 Social issues

Abstract en español

La creciente difusión de las tecnologías digitales y la intensificación de las actividades en línea han provocado un aumento significativo de los fenómenos de cibervictimización, haciendo necesario comprender los factores conductuales y situacionales que influyen en la exposición al riesgo. Esta tesis analiza el papel de la adicción al smartphone, el comportamiento de juego en línea y las rutinas digitales en la explicación de la vulnerabilidad a la cibervictimización, integrando la Lifestyle-Routine Activity Theory (L-RAT) en un marco empírico unificado.

El estudio se basa en un conjunto de datos de carácter nacional y adopta un enfoque metodológico integrado que combina Partial Least Squares Structural Equation Modeling (PLS-SEM) y técnicas de Machine Learning. El PLS-SEM se utiliza para contrastar las relaciones estructurales entre los constructos teóricos, modelando la L-RAT como un constructo formativo de segundo orden, mientras que un modelo Random Forest se emplea para evaluar el rendimiento predictivo y la importancia de las variables.

Los resultados indican que la adicción al smartphone y determinadas dimensiones de la L-RAT, especialmente la Proximity to motivated offenders, se asocian significativamente con el riesgo de cibervictimización. El juego en línea emerge como un factor que amplifica los mecanismos de exposición rutinaria al riesgo, mientras que el apoyo social desempeña un papel protector. La integración de la modelización explicativa y predictiva evidencia una convergencia entre la relevancia teórica y la importancia predictiva de las variables clave.

El estudio aporta una doble contribución: refuerza la aplicabilidad de la Lifestyle-Routine Activity Theory en entornos digitales contemporáneos y demuestra el valor de integrar la modelización estructural con algoritmos de aprendizaje automático para el análisis del riesgo conductual en línea.

Palabras clave:

Cibervictimización; Adicción al smartphone; Juego en línea; Lifestyle-Routine Activity Theory; PLS-SEM; Machine Learning; Random Forest; Conducta digital

Códigos UNESCO:

1203 Computer Science

1203.12 Databases

1209 Statistics

1209.01 Analytical Statistics

1209.03 Data Analysis

1209.13 Statistical Inference Techniques

1209.14 Statistical Prediction Techniques

6114 Social Psychology

6310 Social issues

Contents

Acknowledgments	i
Abstract.....	iii
Abstract in italiano.....	v
Abstract en español.....	vii
Contents	ix
1 Introduction	1
1.1. Master's thesis justification.....	1
1.2. Objectives	2
1.3. Document structure	4
2 Theoretical framework.....	7
2.1. Lifestyle-Routine Activity Theory (LRAT).....	7
2.1.1. Exposure to motivated offenders.....	10
2.1.2. Proximity to motivated offenders.....	11
2.1.3. Target suitability.....	11
2.1.4. Capable guardianship.....	12
2.2. Online gambling behaviour.....	13
2.3. Smartphone Addiction (SA)	15
2.4. Social Digital Pressure (SDP).....	18
2.5. Social Support (SS).....	19
2.6. Conceptual model	21
2.7. Contribution of this study.....	23
2.7.1. Scientific contribution.....	23
2.7.2. Methodological contribution	24
2.7.3. Policy relevance	24
3 Methodology.....	25
3.1. Regression types methods and their limits	26
3.1.1. Examples of regression-type methods	26
3.1.2. Limitations of regression-type methods	27
3.2. Structural Equation Modelling.....	27
3.2.1. SEM: foundations and rationale for the adoption of PLS-SEM.....	27
3.2.2. Components of a PLS-SEM.....	30

3.2.3.	Evaluation of Reflective Measurement Models	35
3.2.4.	Evaluation of Formative Measurement Models	38
3.2.5.	Evaluation of the Structural Model	40
3.2.6.	Higher-Order Construct.....	43
3.2.7.	Mediation analysis	45
3.2.8.	Moderation analysis.....	48
3.2.9.	Summary of PLS-SEM evaluation procedures and criteria	49
3.3.	Machine Learning	50
3.3.1.	Machine Learning and predictive modelling.....	50
3.3.2.	Machine Learning paradigms and supervised learning models.....	52
3.3.3.	Machine Learning lifecycle	54
3.3.4.	Random Forest.....	56
3.3.5.	Evaluation of the Machine Learning model	59
4	Data, variables, and modelling strategy	63
4.1.	Dataset description	63
4.2.	Data preparation and preprocessing.....	64
4.2.1.	Item selection and theoretical alignment.....	64
4.2.2.	Temporal filtering and removal of invalid observations	65
4.2.3.	Recoding of sociodemographic variables	65
4.2.4.	Consistency checks and reverse coding	65
4.2.5.	Data scaling for PLS-SEM and Machine Learning	66
4.2.6.	Descriptive statistics and reliability checks	66
4.3.	Variables, constructs, and measurement scales.....	67
4.3.1.	Dependent variable: Cybervictimization risk	67
4.3.2.	L-RAT framework	68
4.3.3.	Gambling	74
4.3.4.	Smartphone Addiction	75
4.3.5.	Social Digital Pressure	76
4.3.6.	Social Support	77
4.3.7.	Sociodemographic variables.....	78
4.4.	Structural Modelling Strategy: PLS-SEM	80
4.4.1.	Operative Strategy: First-Order Model specification.....	80
4.4.2.	Operative Strategy: Higher-Order Model Specification.....	81
4.5.	Predictive Modelling Strategy: Machine Learning	82
4.5.1.	Operative strategy: Variable definition and feature construction	82
4.5.2.	Operative strategy: Training, validation, and model estimation	83
4.5.3.	Operative strategy: Model evaluation and interpretation	84

5	Results and discussion.....	87
5.1.	PLS-SEM results	87
5.1.1.	Reflective measurement model	87
5.1.2.	Formative measurement model	89
5.1.3.	Structural model	91
5.1.4.	Mediation analysis	95
5.2.	Machine Learning results.....	97
5.2.1.	Model performance evaluation	97
5.2.2.	Variable importance analysis	98
5.3.	Discussion	100
5.3.1.	Convergence of explanatory and predictive evidence.....	100
5.3.2.	L-RAT as central mechanism of cybervictimization	100
5.3.3.	Smartphone addiction and behavioural vulnerability.....	101
5.3.4.	Gambling behaviours and cyber risk	101
5.3.5.	Psychosocial antecedents of Smartphone Addiction	102
5.3.6.	The role of demographic factors	102
5.3.7.	Integrating criminological theory and behavioural addiction	103
6	Conclusions, limitations and future investigation.....	105
6.1.	Conclusions.....	105
6.2.	Policy, practical, and social implications.....	107
6.3.	Limitations and future research.....	108
	Time planning and budgeting.....	111
	Ethical, legal, and professional considerations	115
	Contribution to the sustainable development goals (SDGs).....	117
	References	119
	List of Figures.....	127
	List of Tables	129
	List of Equations.....	131
	List of acronyms and symbols	133
A	Appendix - R code.....	135

1 Introduction

1.1. Master's thesis justification

The pervasive diffusion of digital technologies has profoundly transformed everyday life, reshaping social interaction, economic activities, and leisure practices. Smartphones, online platforms, and digital services have become deeply embedded in daily routines, progressively blurring the boundaries between offline and online environments. Recent global evidence highlights a continuous growth in digital connectivity, with the number of Internet users steadily increasing worldwide (International Telecommunication Union, 2025). At the same time, smartphones have consolidated their role as the dominant device for accessing online services, communication platforms, and digital content, becoming the primary gateway to online environments for large segments of the population (Waweru, 2025).

This process of digital transformation has generated substantial benefits at the individual level, including improved access to information, enhanced efficiency in daily activities, and expanded opportunities for social participation and well-being (Organisation for Economic Co-operation and Development, 2019; Vuorre & Przybylski, 2024). However, the same transformation has also altered individuals' exposure to risk, increasing the frequency, intensity, and diversity of potentially harmful online interactions. As digital technologies become increasingly integrated into everyday routines, individuals are more frequently exposed to cyber threats, online fraud, and other forms of cyber-enabled (Islam & Bhuiyan, 2022; Kramer, 2025).

In this context, cybercrime has emerged as a critical social and economic issue. The expansion of digital interactions has been accompanied by a parallel increase in online fraud, scams, identity theft, and security incidents affecting individual users. Research in this domain shows that cyber-victimization is not randomly distributed across the population, but is systematically associated with behavioural patterns, levels of online exposure, and individual characteristics. In particular, susceptibility to cybercrime has been linked to higher levels of online engagement, risk-taking behaviours, and specific cognitive and emotional profiles (Whitty, 2019). These findings challenge purely technical interpretations of cybercrime and point to the need for behavioural and psychosocial explanations of digital vulnerability.

Among the digital activities that may intensify exposure to online risks, online gambling represents a particularly relevant case. Online gambling platforms combine

continuous Internet connectivity, financial transactions, and emotionally charged decision-making processes, often mediated through smartphones and personal digital devices. Empirical evidence indicates that higher gambling frequency is associated with impulsivity, sensation seeking, and problematic patterns of technology use, especially among adolescents and young adults (Testa et al., 2025). Moreover, digital gambling environments frequently rely on persuasive design features, targeted advertising, and rapid feedback mechanisms that may impair users' risk perception and self-regulation, potentially increasing vulnerability to harmful online outcomes.

From a broader theoretical perspective, the relationship between digitalization and risk can be interpreted as the outcome of a structural transformation in how individuals interact with digital technologies. Recent contributions in the field of digital society research highlight that digital transformation does not merely introduce new tools, but reshapes everyday practices, social interactions, and exposure to both opportunities and threats (Islam & Bhuiyan, 2022; Organisation for Economic Co-operation and Development, 2019). Within this framework, online risks are not solely the result of malicious actors, but emerge from routine digital activities embedded in daily life.

Despite the growing literature on cybercrime and online gambling, these streams of research have largely developed independently. Studies on cybercrime victimization often focus on individual security behaviours and awareness (Anderson & Agarwal, 2010), while gambling research has traditionally emphasised clinical or pathological interpretations of problematic play. More recent contributions have reframed gambling as a broader public health and behavioural issue, highlighting how digital environments and online platforms can amplify exposure to harm beyond strictly clinical dimensions (Wardle et al., 2019). From this perspective, online gambling can be understood as a risky digital behaviour embedded in everyday online routines, characterised by high accessibility, sustained exposure, and frequent interaction with digital systems. Such routine engagement with online environments has been shown to increase individuals' vulnerability to a range of online risks, as patterns of online activity and exposure play a central role in shaping the likelihood of cyber-victimization (Holt & Bossler, 2015).

This master's thesis is therefore justified by the need to bridge these fragmented approaches through an integrated analytical perspective. By combining psychosocial factors, online gambling behaviours, and a digital adaptation of Lifestyle-Routine Activity Theory, the study aims to offer a more comprehensive explanation of cyber-victimization risk in contemporary digital societies.

1.2. Objectives

This master's thesis is developed within the research projects "Smartphone addiction, loss of psychosocial well-being and online victimization of online gamblers

(CIBERADICJUEGO2024), Reference SUBV24/00012” and “The commodification of human behaviour: Data sharing and privacy threats among online gamblers, Reference SUBV25/00014” funded and supported by the Spanish Ministry of Social Rights, Consumption and Agenda 2030, DG Gambling. Alberto Urueña, Supervisor of this work, is IP of both projects.

The primary objective of this study is to analyse how online gambling behaviour, smartphone addiction, and digital routines influence the likelihood of experiencing cyber-victimization. Grounded in the Lifestyle-Routine Activity Theory, the research adopts a behavioural and psychosocial perspective to explain patterns of digital risk and vulnerability using data drawn from a national cybersecurity survey. In doing so, the thesis aims to move beyond purely technical explanations of cybercrime by integrating individual-level behavioural factors with situational exposure mechanisms.

More specifically, the objectives of the present study are the following:

- To conceptually examine online gambling behaviour, smartphone addiction, and Lifestyle-Routine Activity Theory, assessing their relevance as complementary frameworks for explaining cybercrime victimization.
- To analyse the relationships between psychosocial variables, risky digital behaviours, and online victimization, identifying the mechanisms through which certain user profiles become more exposed to cyber risks.
- To apply Partial Least Squares Structural Equation Modelling (PLS-SEM) as a methodological approach suitable for behavioural and exploratory research, in order to estimate complex relationships among latent constructs.
- To develop and compare predictive models based on PLS-SEM and Machine Learning techniques, evaluating their respective explanatory and predictive performance in the context of cyber-victimization risk.
- To discuss the theoretical and practical implications of the results, providing insights that may support future prevention strategies and public policies aimed at reducing cybercrime victimization and mitigating harmful digital behaviours.

Through these objectives, the thesis seeks to contribute both to the academic literature on digital risk behaviours and to the development of evidence-based approaches for understanding and preventing cyber-victimization in increasingly digitised societies.

1.3. Document structure

This thesis is organised into four main chapters, each addressing a specific stage of the research process and contributing to the overall objective of understanding cybervictimization risk in digital environments.

Chapter 1 introduces the study by presenting its context, motivation, and objectives. The chapter situates the research within the broader framework of the projects “Smartphone addiction, loss of psychosocial well-being and online victimization of online gamblers (*CIBERADICJUEGO2024*), Reference SUBV24/00012” and “The commodification of human behaviour: Data sharing and privacy threats among online gamblers, Reference SUBV25/00014”. It outlines the public and scientific relevance of the topic and identifies the research gap concerning smartphone addiction, digital routines, and cybervictimization. It also defines the theoretical, methodological, and policy-oriented objectives guiding the study and provides an overview of the document structure.

Chapter 2 develops the theoretical framework and literature review underpinning the research. It examines the main psychosocial and behavioural constructs considered in the study, including online gambling behaviour, smartphone addiction, social digital pressure, and social support. The chapter then introduces the Lifestyle-Routine Activity Theory and its adaptation to digital environments, justifying the specification of a second-order formative construct. Based on the reviewed literature, the theoretical model and research hypotheses are formulated, and the expected scientific, methodological, and policy contributions of the study are discussed.

Chapter 3 presents the methodological framework of the study. It describes the research design, the characteristics of the dataset, and the statistical techniques adopted to address the research questions. The chapter details the specification of the PLS-SEM model, including the measurement and structural components, the modelling of the Lifestyle-Routine Activity Theory as a higher-order construct, and the criteria used to assess model validity and predictive relevance. It also introduces the machine learning approach and justifies its use as a complementary predictive tool.

Chapter 4 focuses on the empirical setup of the analysis. It provides a detailed description of the data source and sample, outlines the procedures used for data cleaning and preprocessing, and explains the operationalization of all variables included in the study. The chapter further describes the modelling strategies adopted for both structural equation modelling and machine learning, highlighting their distinct objectives and methodological assumptions.

Chapter 5 presents and critically discusses the empirical results of the study. It first reports the outcomes of the PLS-SEM analysis, including the assessment of the measurement model, the evaluation of structural relationships, and the analysis of

explanatory power and mediation effects. It then introduces the results of the Machine Learning model, focusing on predictive performance and variable importance. The chapter concludes with an integrated interpretation of explanatory and predictive findings, highlighting the complementarity between structural modelling and data-driven approaches.

Finally, Chapter 6 provides the overall conclusions of the study. It synthesizes the main theoretical and empirical contributions, discusses their policy and practical implications, and reflects on the broader relevance of the findings. The chapter also addresses the main methodological and analytical limitations of the research and outlines directions for future investigation. Together, Chapter 5 and 6 consolidate the explanatory and predictive evidence generated by the study and position the research within the existing literature.

2 Theoretical framework

This chapter develops the theoretical framework underpinning the present study. Building on criminological, behavioural, and psychosocial literature, it integrates multiple strands of research to explain cybervictimization as the outcome of interconnected digital practices rather than isolated technical failures or individual vulnerabilities. The chapter moves beyond single-factor explanations by adopting a multi-layered perspective that combines behavioural addictions, social dynamics, and opportunity-based criminological theory.

Specifically, the framework brings together Lifestyle-Routine Activity Theory, online gambling behaviours, smartphone addiction, social digital pressure, and social support into a unified conceptual model. Each construct is introduced and theoretically grounded, with particular attention paid to how digital environments reshape exposure to risk, social expectations, and behavioural regulation. Rather than treating these factors as independent predictors, the chapter emphasises their interdependence and cumulative effects on cybervictimization risk.

By progressively introducing hypotheses across sections, the chapter establishes a coherent pathway from theory to empirical testing. The final sections consolidate these relationships into an integrated theoretical model, clarify the presence of control variables, and outline the scientific, methodological, and policy contributions of the study.

2.1. Lifestyle-Routine Activity Theory (LRAT)

Research on crime victimization has progressively moved away from strictly offender-centred explanations, toward analytical frameworks that explicitly incorporate victims' everyday activities and the social contexts in which these activities take place. Early criminological research primarily focused on individual motivations, psychological traits, or broad structural determinants such as social disadvantage, deviance, and inequality. Within these perspectives, victimization was often treated as a secondary or incidental outcome of criminal behaviour, receiving limited analytical attention as an object of explanation in its own right (Rock, 2002; Walklate, 2011). As a result, victims were frequently conceptualized as passive actors, with little emphasis placed on how their daily routines or situational contexts might influence their exposure to risk.

From the late twentieth century onward, however, a growing body of empirical research began to challenge this assumption. Studies in victimology consistently demonstrated that victimization is not evenly or randomly distributed across individuals or social groups, but instead follows systematic and predictable patterns (Hindelang et al., 1978; Miethe & Meier, 1994). These patterns were shown to be closely linked to how individuals structure their daily lives, including where they spend their time, the activities they routinely engage in, and the social environments they inhabit. This empirical shift redirected scholarly attention toward the role of lifestyles and routine behaviours as key mechanisms shaping exposure to risky situations and opportunities for crime. Within this context, theoretical approaches integrating situational opportunity structures with everyday social practices emerged as particularly well suited to explaining victimization dynamics, laying the conceptual foundations for the development of Lifestyle-Routine Activity Theory.

Lifestyle-Routine Activity Theory (LRAT) represents one of the most influential theoretical frameworks for explaining crime victimization. LRAT originates from the integration of two complementary approaches: *Routine Activity Theory*, introduced by Cohen and Felson (1979), and the *Lifestyle model of victimization*, developed by Hindelang, Gottfredson, and Garofalo (1978). Both frameworks emerged in response to empirical observations of rising property crime in the United States during the 1970s and sought to explain crime trends without relying exclusively on offender motivation or macro-social variables.

Routine Activity Theory posits that criminal events occur when three elements converge in space and time: a motivated offender, a suitable target, and the absence of capable guardianship (Cohen & Felson, 1979). This approach highlights the situational and environmental nature of crime opportunities, emphasizing that changes in everyday routines, such as increased mobility or time spent outside the home, can systematically affect crime rates. The *Lifestyle model* complements this view by focusing on how individuals' routine activities, both occupational and leisure-related, shape their exposure to risky contexts and interactions, thereby influencing victimization risk (Hindelang et al., 1978; Garofalo, 1986).

The integration of these perspectives led to LRAT, which conceptualizes victimization risk as a function of how lifestyles and routine activities structure individuals' exposure to motivated offenders, proximity to risky situations, target suitability, and levels of guardianship. From this standpoint, victimization is not merely the result of offender behaviour, but the outcome of patterned social activities embedded in everyday life.

While LRAT was initially developed to explain victimization in physical and social environments such as urban settings and everyday movements beyond the home, its fundamental logic is equally applicable in digital ecosystems where crime opportunities have shifted. The expansion of online activities and the pervasive

diffusion of internet-connected devices have profoundly altered how individuals interact, socialize, and conduct economic transactions, thereby reshaping opportunities for crime.

In this context, cybercrime refers to illegal behaviours that exploit digital technologies and networks to cause harm to individuals, organisations, and public institutions. According to the Osservatori Digital Innovation of the Politecnico di Milano, cybercrime encompasses any act committed through the misuse of information systems or connected devices with the intent to defraud, steal data, disrupt services, or otherwise compromise digital assets. It includes a range of behaviours from data breaches and ransomware attacks to fraud and identity theft, reflecting both intentional exploitations and opportunistic attacks facilitated by digital interconnectivity (Piva, 2019).

Recent national data illustrate the scale and growth of cybercrime in contemporary digital societies. In 2024, globally there were 3,541 publicly reported severe cyber incidents, of which approximately 10 per cent occurred in Italy, with large organisations particularly affected; 73 per cent of larger Italian enterprises reported suffering cyber-attacks in the same year. At the same time, the national cybersecurity market expanded to an estimated 2.48 billion euros in investments, signalling both the growing threat landscape and increasing attention to defensive strategies (Piva, 2019). For comparative context, reports from Spain indicated a significant rise in cybercrime incidents during the same period, with estimates suggesting over 237,000 cybercrime cases reported in the first seven months of 2024, driven by increases in ransomware, fraud, and online scams (*Cybersecurity in Spain: Challenges and Opportunities*, n.d.). These figures underline the degree to which digital victimization has become a systemic concern across different European jurisdictions.

Early empirical research on computer crime victimization had already challenged purely technical explanations of cybercrime, showing that individual routines and opportunity structures play a central role in shaping victimization risk. In particular, Choi (2008) provided one of the first empirical assessments integrating routine activity and lifestyle perspectives into the study of computer crime victimization, demonstrating that online routines and protective behaviours significantly affect individuals' likelihood of being victimized. Crucially, research consistently highlights that cybervictimization cannot be explained solely by technological vulnerabilities of systems and networks. Instead, human behaviours, routine practices, and levels of cybersecurity awareness play a central role in shaping exposure to digital risk. Empirical studies show that individuals' routinely online activities, decision-making processes, and behavioural patterns interact with technological safeguards to influence vulnerability to cyber threats (Leukfeldt & Yar, 2016; Whitty, 2019; Herrero, Torres, Vivas, Hidalgo, et al., 2021).

From a behavioural perspective, factors such as frequency of online engagement, types of platforms used, and attitudes toward privacy and security contribute to differentiated levels of cyber risk. In this context, the literature increasingly emphasises the relevance of the “human factor”, understood as the role of user behaviour, awareness, and cognitive biases in facilitating or mitigating cybercrime victimization. This perspective aligns with institutional evidence highlighting how human actions often represent a critical vulnerability within digital security systems (Hadlington, 2017; Akdemir & Lawless, 2020; Piva, 2019).

For these reasons, several scholars have adapted Lifestyle-Routine Activity Theory to the analysis of cybervictimization, demonstrating that its core theoretical structure remains coherent and applicable when key constructs are reinterpreted for digital environments. In cyberspace, physical co-presence is replaced by interaction within digital platforms, while spatial convergence is reconceptualised as situational convergence within online communication and transactional contexts. Despite the absence of geographical proximity, empirical research shows that mechanisms such as exposure, proximity, target suitability, and guardianship continue to shape victimization risk in meaningful and systematic ways (Leukfeldt & Yar, 2016; Miró-Llinares et al., 2020; Herrero, Torres, Vivas, Hidalgo, et al., 2021).

Consequently, LRAT provides a valid conceptual foundation for understanding cybervictimization as the outcome of interconnected digital practices embedded in users’ routines, rather than isolated technological failures or random events.

2.1.1. Exposure to motivated offenders

Within the LRAT framework, exposure to motivated offenders refers to the extent to which individuals’ routine activities systematically place them in contexts where potential offenders are likely to operate. In physical environments, exposure is traditionally associated with time spent in public spaces, unstructured activities, or routine patterns that increase the likelihood of encountering offenders, even in the absence of direct interaction (Cohen & Felson, 1979; Garofalo, 1986).

In digital environments, exposure must be reinterpreted in terms of online presence and activity intensity rather than physical co-presence. Exposure reflects how frequently and extensively individuals engage with online environments that may harbour malicious actors, such as open platforms, unregulated websites, or services characterised by weak security controls. Higher exposure does not imply intentional interaction with offenders, but rather an increased probability of encountering threats as a by-product of routine digital behaviour (Herrero, Torres, Vivas, Hidalgo, et al., 2021).

Empirical research highlights that exposure in cyberspace is shaped by both the frequency and type of online activities performed. Activities such as downloading software from unofficial sources, accessing free content platforms, or engaging with

services that require extensive data sharing have been associated with higher cybervictimization risk (Leukfeldt & Yar, 2016; Miró-Llinares et al., 2020). From this perspective, exposure captures a structural dimension of digital lifestyles: individuals who spend more time online and across a wider range of services are statistically more likely to be present in environments where motivated offenders operate.

2.1.2. Proximity to motivated offenders

Proximity to motivated offenders traditionally refers to the convergence of offenders and potential victims in space and time, which increases the opportunity for crime to occur (Cohen & Felson, 1979). In physical contexts, proximity is influenced by geographical location, time of day, and patterns of movement that facilitate encounters between offenders and targets.

In cyberspace, physical distance becomes irrelevant, requiring a conceptual shift toward situational and interactional proximity. Proximity in digital environments reflects how closely and how often individuals come into contact with potential offenders through online interactions, communication channels, or digital transactions. This includes both intentional interactions (e.g.: online marketplaces, gaming platforms) and unintentional contacts (e.g.: receiving phishing emails, fraudulent messages, or malicious links).

Research suggests that digital proximity is shaped by users' integration into online networks and communication ecosystems. Frequent use of instant messaging services, social networks, and interactive platforms increases the likelihood of encountering malicious actors, regardless of geographical location (Akdemir & Lawless, 2020; Herrero, Torres, Vivas, Hidalgo, et al., 2021). In this sense, proximity captures the density and immediacy of digital interactions, which may facilitate rapid targeting and exploitation by offenders.

Importantly, proximity does not require awareness on the victim side. Individuals may be proximate to offenders simply by participating in highly interconnected digital environments where malicious actors operate alongside legitimate users. This reinforces the relevance of proximity as a situational risk factor within digital routine activity patterns.

2.1.3. Target suitability

Target suitability refers to the degree to which a potential victim appears attractive, accessible, or vulnerable from the offender's perspective. In traditional crime settings, suitability is influenced by visible cues such as physical vulnerability, isolation, or possession of valuable goods (Hindelang et al., 1978).

In digital environments, target suitability is constructed through users' online profiles, behaviours, and information disclosure practices. Rather than physical attributes, suitability is signalled through digital traces, such as the amount of personal

information shared, behavioural consistency, and engagement in risky online practices. Offenders rely on these signals to identify individuals who are more likely to respond to fraudulent attempts or lack adequate protective measures.

Empirical research shows that frequent engagement in transactional online activities and routine exposure to digital services significantly increase the likelihood of being targeted in cyberspace (Pratt et al., 2010; Herrero, Torres, Vivas, Hidalgo, et al., 2021). Importantly, time spent online alone is insufficient to explain target suitability. Rather, the qualitative nature of online activities plays a crucial role. Users who regularly participate in digital environments involving registration procedures, payment information, or identity verification may inadvertently increase their visibility and perceived attractiveness as targets, particularly when such activities are combined with risky behavioural patterns or limited protective practices (Macaulay et al., 2020).

Target suitability thus represents a behavioural and informational dimension of vulnerability, capturing how individuals' routine digital practices shape offenders' perceptions of potential success.

2.1.4. Capable guardianship

Capable guardianship refers to the presence of protective mechanisms that can prevent, deter, or mitigate criminal acts. In the original LRAT formulation, guardianship includes both formal actors (e.g.: law enforcement) and informal forms of protection, such as bystanders or self-protective behaviours by potential victims (Cohen & Felson, 1979).

In digital contexts, guardianship is predominantly enacted at the individual and technological level, as most online activities occur in private or semi-private settings. Digital guardianship encompasses a combination of behavioural strategies and technical safeguards that reduce exposure to cyber threats. Scholars commonly distinguish between active guardianship, involving cautious behaviours and informed decision-making, and passive guardianship, consisting of automated or system-level protections (Anderson & Agarwal, 2010).

Active guardianship includes behaviours such as verifying sources, avoiding suspicious links, and refraining from downloading unverified content. Passive guardianship involves the use of security tools such as antivirus software, firewalls, authentication mechanisms, and operating system protections. While technological safeguards provide a baseline level of protection, their effectiveness often depends on users' awareness and correct usage.

Research consistently shows that insufficient guardianship, either due to lack of awareness or overreliance on technology, significantly increases cybervictimization risk (Hadlington, 2017; Whitty, 2019). Consequently, guardianship represents a critical

moderating component within the LRAT framework, capable of attenuating the effects of exposure, proximity, and target suitability on cybercrime victimization.

Although exposure, proximity, target suitability, and guardianship represent conceptually distinct dimensions, Lifestyle-Routine Activity Theory conceptualizes them as interrelated components of a broader behavioural framework, in which victimization risk emerges from the configuration and cumulative interaction of routine activities rather than from isolated factors. In line with this theoretical perspective, empirical research increasingly conceptualizes and analyses these dimensions jointly, showing that their combined effect plays a central role in shaping cybervictimization risk in digital environments and highlighting the cumulative nature of routine activities (Herrero, Torres, Vivas, Hidalgo, et al., 2021; Nahar & Bhuyain, 2021).

From a methodological perspective, modelling LRAT as a second-order formative construct is consistent with the logic of composite-based structural equation modelling, as the dimensions jointly define the construct rather than reflecting a single underlying trait (Diamantopoulos et al., 2008; Hair et al., 2021). This approach allows the theoretical integrity of LRAT to be preserved while empirically assessing the distinct contribution of each component.

Based on the existing literature, LRAT is expected to play a central role in explaining cybervictimization risk in digital environments. Accordingly, the first hypothesis of this study is formulated as follows:

H1: *Lifestyle-Routine Activity patterns are positively associated with cybervictimization risk, such that higher levels of exposure, proximity, target suitability, and lower levels of guardianship increase the likelihood of cybervictimization.*

2.2. Online gambling behaviour

Online gambling has become an increasingly widespread form of digital leisure as a consequence of the diffusion of internet-connected devices and the normalization of online financial transactions. Gambling activities, traditionally confined to physical venues and subject to spatial and temporal constraints, have progressively migrated to online platforms, allowing individuals to place bets or participate in games of chance at any time and from virtually any location. This transformation has significantly altered the accessibility, frequency, and intensity of gambling behaviours, raising concerns about their psychosocial and behavioural consequences (Gainsbury et al., 2015; Wardle et al., 2019).

Gambling can be defined as the wagering of money or other valuables on the outcome of events whose results are at least partly determined by chance. In online environments, this definition encompasses activities such as sports betting, online casino games, poker, and other digital gambling formats based on probabilistic

outcomes rather than skill alone. While gambling is legally regulated and socially tolerated in many countries, its online expansion has amplified exposure to risk by lowering entry barriers, increasing immediacy of rewards, and facilitating continuous engagement. As a result, online gambling has been increasingly associated with problematic behaviours, financial difficulties, and adverse mental health outcomes (Stefanovics & Potenza, 2022; Wardle et al., 2019).

A substantial body of empirical research indicates that online gambling is associated with a higher risk profile compared to traditional offline gambling. Studies consistently show that online gamblers display elevated levels of impulsivity, sensation seeking, and reduced self-control, psychological traits that are strongly linked to the development of problematic gambling behaviours (Kairouz et al., 2012; Testa et al., 2025). These individual vulnerabilities interact with the structural characteristics of online gambling platforms, such as rapid betting cycles, instant feedback, and algorithmically personalised incentives, fostering patterns of use comparable to those observed in highly addictive gambling modalities, including electronic gaming machines (Gainsbury, 2015).

From a clinical and public health perspective, excessive and uncontrolled gambling has been recognised as a behavioural addiction, commonly referred to as gambling disorder. This condition is characterised by impaired control over gambling behaviour, persistence despite negative consequences, and significant personal, social, and economic harm (Stefanovics & Potenza, 2022). However, recent literature has increasingly criticised approaches that frame gambling-related harm exclusively in terms of individual pathology, arguing that such perspectives underestimate the role of structural and commercial factors embedded in online gambling environments. In particular, critiques of the “responsible gambling” paradigm highlight how platform design, persuasive technologies, and profit-driven incentives systematically contribute to risky patterns of use (Wardle et al., 2019; Livingstone & Rintoul, 2020).

Beyond gambling-specific outcomes, online gambling is particularly relevant for the study of cybervictimization due to its association with risky online behaviours and intensified digital engagement. Empirical studies show that individuals who engage in online gambling are more likely to spend extended periods online, perform frequent financial transactions, and disclose sensitive personal and payment information (Canale et al., 2016; Escario & Wilkinson, 2020). These practices place users in digital environments characterised by heightened exposure to cyber threats, increasing the likelihood of encountering malicious actors and experiencing online victimization. From this perspective, online gambling may exert a direct effect on cybervictimization risk by fostering behavioural patterns associated with reduced caution and increased vulnerability.

In addition to this direct relationship, online gambling is likely to influence cybervictimization indirectly by reshaping individuals’ lifestyle and routine online

activities. Participation in online gambling platforms requires repeated interactions with complex digital ecosystems involving financial services, authentication procedures, and data exchanges across multiple platforms and devices. These routine practices may increase exposure to motivated offenders, amplify target suitability through repeated financial activity, and weaken guardianship by fostering overconfidence or habituation to digital risk (Gainsbury et al., 2015; Holt & Bossler, 2015). Consequently, online gambling can be conceptualized as a behavioural condition that modifies the mechanisms described by Lifestyle-Routine Activity Theory, indirectly elevating cybervictimization risk through its impact on exposure, proximity, target suitability, and guardianship.

In the present study, online gambling behaviour is therefore treated as a relevant behavioural factor influencing cybervictimization through both direct and indirect pathways. Specifically, two hypotheses are formulated to capture these mechanisms:

H2: *The presence of online gambling behaviours is positively associated with cybervictimization risk, such that individuals who engage in online gambling exhibit higher levels of cybervictimization.*

H3: *Online gambling behaviours indirectly increase cybervictimization risk by shaping individuals' lifestyle and routine activity patterns, operating through the L-RAT framework as a mediating mechanism.*

2.3. Smartphone Addiction (SA)

Smartphones have become one of the most pervasive digital technologies of contemporary societies, deeply embedded in everyday communication, information seeking, entertainment, and access to online services. Since their rapid diffusion following the introduction of touchscreen smartphones in the late 2000s, these devices have evolved from simple communication tools into multifunctional platforms that integrate social networking, entertainment, financial transactions, and constant internet connectivity. As a consequence, concerns have increasingly emerged regarding patterns of excessive or poorly regulated smartphone use and their potential implications for individual well-being and online vulnerability (Panova & Carbonell, 2018; Busch & McCarthy, 2021).

Within the scientific literature, the term Smartphone Addiction (SA), often used interchangeably with problematic smartphone use, has been introduced to describe patterns of smartphone engagement characterised by loss of behavioural control, excessive time investment, and negative consequences across psychological, social, or functional domains. Importantly, this construct remains the subject of ongoing debate. Smartphone addiction is not formally recognised as a distinct clinical disorder in the Diagnostic and Statistical Manual of Mental Disorders (DSM-5-TR) (American Psychiatric Association, 2022), and there is no universally accepted diagnostic

threshold distinguishing normative from pathological use. Recent reviews highlight that most empirical studies rely on self-report instruments applied to non-clinical samples, often extrapolating criteria from substance-related or behavioural addictions, which raises concerns about over-pathologisation and construct validity (Ting & Chen, 2020; James et al., 2023).

Rather than framing smartphone addiction as a clinical condition, contemporary research increasingly conceptualises it as a behavioural pattern reflecting difficulties in self-regulation within digitally saturated environments. From this perspective, problematic smartphone use is understood as the outcome of repeated reinforcement mechanisms, social expectations of constant availability, and contextual pressures that encourage persistent connectivity. Empirical evidence suggests that specific behavioural indicators, such as prolonged daily use, rapid engagement after waking, and difficulty disengaging from the device, are more informative than simple frequency measures when assessing problematic use patterns (Haug et al., 2015; Ting & Chen, 2020).

A substantial body of research indicates that smartphone addiction does not develop in isolation, but emerges from the interaction of individual, social, and contextual factors. Psychosocial variables such as fear of missing out, social pressure, emotional distress, and deficits in perceived social support have been consistently linked to higher levels of problematic smartphone use (Beyens et al., 2016; Busch & McCarthy, 2021). Longitudinal evidence further shows that elevated smartphone addiction predicts subsequent declines in social support, suggesting a reinforcing cycle in which excessive smartphone use may erode offline social resources over time (Herrero, Torres, et al., 2019; Herrero, Urueña, et al., 2019b). These findings support an ecological interpretation of smartphone addiction, whereby individual behaviours are shaped by broader social and environmental contexts rather than by isolated personal traits.

Beyond its psychosocial correlates, smartphone addiction is particularly relevant for the study of cybervictimization due to its implications for users' digital routines and online behaviours. Smartphones function as the primary access point to online environments for many individuals, facilitating continuous connectivity, multitasking across platforms, and frequent interaction with digital services. Research has shown that higher levels of smartphone addiction are associated with reduced caution in online interactions, greater engagement in risky digital behaviours, and increased exposure to technology-related harms (Herrero, Urueña, et al., 2019a; Hadlington, 2017).

Consistent empirical evidence indicates a direct association between smartphone addiction and cybercrime victimization. In particular, (Herrero, Torres, Vivas, Hidalgo, et al., 2021) demonstrated that individuals with higher levels of smartphone addiction are significantly more likely to experience cybervictimization, even when

controlling for other psychological and behavioural factors. Excessive smartphone use may directly elevate cyber risk by increasing time spent online, intensifying situational exposure, and impairing users' capacity to recognise or avoid potentially harmful digital contexts. From this standpoint, smartphone addiction can be conceptualised as a direct vulnerability factor within the digital ecosystem.

At the same time, smartphone addiction is likely to influence cybervictimization indirectly by reshaping individuals' lifestyle and routine online activities. Excessive reliance on smartphones modifies daily patterns of connectivity, increases engagement with multiple online platforms, and amplifies exposure to environments in which motivated offenders may operate. These behavioural changes may also weaken effective guardianship through habituation to digital risk, overconfidence in personal skills, or reduced self-regulatory capacity. Such mechanisms closely align with the core components of Lifestyle-Routine Activity Theory, suggesting that smartphone addiction may indirectly increase cybervictimization risk through its impact on exposure, proximity, target suitability, and guardianship (Herrero, Torres, Vivas, Hidalgo, et al., 2021; Akdemir & Lawless, 2020).

Furthermore, emerging evidence suggests that smartphone addiction is associated with other risky digital behaviours, including online gambling. Smartphones provide constant and immediate access to gambling platforms, lowering situational barriers to participation and facilitating impulsive engagement. Studies indicate that individuals with higher levels of smartphone addiction are more likely to engage in online gambling behaviours and display related risk profiles, pointing to an additional indirect pathway through which smartphone addiction may contribute to cybervictimization (Herrero, Torres, Vivas, Hidalgo, et al., 2021; Testa et al., 2025).

Taken together, these findings support a conceptualisation of smartphone addiction as a multifaceted behavioural vulnerability that affects cybervictimization through both direct and indirect mechanisms. In the present study, smartphone addiction is therefore integrated into the theoretical model as a key behavioural construct, leading to the formulation of the following hypotheses:

- H4:** *Smartphone addiction (SA) is positively associated with cybervictimization risk, such that higher levels of SA predict higher levels of cybervictimization.*
- H5:** *Smartphone addiction (SA) indirectly increases cybervictimization risk by shaping individuals' lifestyle and routine activity patterns, operating through the Lifestyle-Routine Activity Theory framework as a mediating mechanism.*
- H6:** *Smartphone addiction (SA) indirectly increases cybervictimization risk by fostering online gambling behaviours, which act as a mediating pathway linking SA to cybervictimization.*

2.4. Social Digital Pressure (SDP)

The diffusion of digital communication technologies has not only expanded opportunities for social interaction, but has also transformed the normative expectations governing availability, responsiveness, and participation in everyday social life. In contemporary digital societies, individuals are increasingly embedded in environments where connectivity is continuous, visibility is persistent, and social interaction is mediated by platforms that implicitly reward immediacy and constant engagement. Within this context, pressure to remain connected does not originate solely from individual preferences, but emerges from the social and technological structures that shape digital interaction.

A growing body of research has highlighted how digital environments generate psychological and social pressure by imposing expectations regarding responsiveness, participation, and online presence (Reer et al., 2019). This pressure is particularly salient in social media and messaging platforms, where indicators such as “last seen,” read receipts, notifications, and engagement metrics make users’ availability and responsiveness observable to others. As a result, remaining connected becomes a socially regulated behaviour rather than a purely voluntary choice.

Social Digital Pressure captures the extent to which digital environments generate implicit and explicit expectations of availability, responsiveness, and continuous participation, particularly through social media and instant messaging platforms. Rather than reflecting an internal psychological disposition, SDP is rooted in contextual and relational demands imposed by digitally mediated social norms. Individuals experiencing high levels of SDP perceive a social obligation to remain connected, to respond promptly, and to maintain visibility within digital networks, even when such engagement is not intrinsically motivated (Gui & Büchi, 2021; Herrero, Torres, Vivas, Arenas, et al., 2021).

This pressure is structurally reinforced by platform-specific design features such as read receipts, activity indicators, notifications, and visibility metrics, which make users’ availability observable and socially accountable. As a result, optional digital engagement is gradually transformed into a perceived social requirement. Empirical research has shown that these dynamics contribute to heightened psychological strain, communication overload, and difficulties in disengaging from digital interactions (Reer et al., 2019).

While related concepts such as the Fear of Missing Out (FoMO) describe individual emotional responses to social exclusion or comparison, SDP operates at a broader contextual level. Studies on FoMO demonstrate that concerns about missing social information or opportunities are associated with more frequent checking behaviours and prolonged engagement with digital platforms (Dhir et al., 2018; Beyens et al., 2016). However, SDP extends beyond individual fear or anxiety by capturing

how social norms and platform affordances jointly transform optional engagement into perceived obligation.

A growing body of longitudinal evidence suggests that SDP represents a relatively stable environmental condition rather than a transient state. Stević (2024) demonstrated that sustained exposure to digitally mediated social expectations predicts increases in compulsive smartphone use and psychological distress over time. Importantly, these effects remain significant after controlling for baseline well-being and usage levels, indicating that digital pressure exerts cumulative influences on behaviour.

Further reinforcing this perspective, Erdem et al. (2025) showed that prolonged exposure to digital social pressure leads not only to higher psychological distress but also to progressive declines in perceived social support over time. These findings suggest that SDP may erode protective social resources, reinforcing reliance on digital interaction and increasing vulnerability to maladaptive technology use.

Recent empirical work has explicitly identified Social Digital Pressure as a meaningful predictor of problematic smartphone use. (Herrero, Torres, Vivas, Arenas, et al., 2021) demonstrated that perceived digital social pressure is significantly associated with higher levels of smartphone addiction, even when controlling for individual traits and psychosocial factors.

Building on this evidence, SDP can be conceptualized as an environmental condition that shapes smartphone use beyond individual motivations. Persistent social expectations to remain available may lead individuals to increase device usage not for intrinsic enjoyment, but to comply with perceived social demands. In this sense, Social Digital Pressure represents a contextual driver linking digitally mediated social environments to individual behavioural outcomes

Based on this theoretical and empirical evidence, Social Digital Pressure is included in the present model as a key antecedent of Smartphone Addiction. Accordingly, the following hypothesis is formulated:

H7: *Higher levels of Social Digital Pressure are positively associated with Smartphone Addiction.*

2.5. Social Support (SS)

Social support has long been recognised as a central protective factor for psychological well-being and adaptive functioning. Broadly defined, social support refers to the perception or experience of being cared for, valued, and embedded within a network of meaningful social relationships, including family members, friends, and significant others. Classic sociological and psychological theories emphasise that social support contributes to resilience by providing emotional reassurance, practical

assistance, and cognitive resources that help individuals cope with stressors and adverse life events (Cohen & Wills, 1985; Thoits, 2011).

The relevance of social support extends beyond general well-being to the domain of digital behaviour. Prior research consistently shows that individuals with lower levels of perceived offline social support are more likely to engage in compensatory behaviours, seeking connection, validation, or emotional regulation through digital technologies. In this context, smartphones may serve as readily available tools for fulfilling unmet social needs, particularly when offline relational environments are perceived as unsatisfactory or insufficient (Herrero, Torres, et al., 2019; Busch & McCarthy, 2021). This compensatory mechanism increases the salience of digital feedback and may foster excessive or poorly regulated smartphone use (Zhou & Cheng, 2025).

Empirical studies consistently show that lower perceived social support is associated with higher levels of smartphone addiction. Cross-sectional and longitudinal evidence indicates that individuals with weaker social ties are more likely to develop problematic patterns of smartphone use, while improvements in social support over time are associated with reductions in addictive behaviours (Herrero, Torres, et al., 2019; Herrero, Urueña, et al., 2019b). These findings support the interpretation of smartphone addiction as, at least in part, a socially embedded phenomenon rather than a purely individual pathology.

Digital environments may partially compensate for deficits in offline support through mechanisms such as likes, comments, and instant messaging. However, these forms of mediated feedback often provide short-term validation rather than stable relational resources. Repetitive engagement with such feedback mechanisms, facilitated by infinite scrolling and algorithmically curated content, can reinforce habitual checking behaviours and contribute to dependence on the smartphone as a primary source of social gratification.

Recent longitudinal research further suggests that social support and digital pressure are dynamically intertwined. Erdem et al. (2025) found that high levels of digital social pressure predict subsequent declines in perceived social support, indicating that demanding digital environments may weaken offline relational resources over time. This erosion of social support may, in turn, increase susceptibility to smartphone addiction by reducing access to alternative coping mechanisms and sources of emotional regulation.

Empirical studies further suggest that insufficient social support is associated with riskier digital trajectories and increased vulnerability to negative online outcomes. In particular, Herrero et al. (2022) showed that lower levels of perceived social support contribute to higher smartphone addiction over time, which in turn increases the likelihood of cybercrime victimization. These findings reinforce the idea

that social resources play a protective role in digital environments, indirectly reducing exposure to online risks.

Taken together, these findings highlight social support as a critical antecedent of smartphone addiction. Lower levels of perceived support not only increase reliance on digital interaction but also reduce resilience against the pressures embedded in digital communication environments. On this basis, the following hypothesis is proposed:

H8: *Lower levels of perceived Social Support are negatively associated with Smartphone Addiction.*

2.6. Conceptual model

This study proposes an integrated theoretical model aimed at explaining cybervictimization risk as the outcome of a set of interrelated social, behavioural, and situational mechanisms operating within contemporary digital environments. Building on the Lifestyle-Routine Activity Theory (LRAT), the model combines criminological theory with insights from research on problematic digital behaviours, social pressure, and psychosocial vulnerability. Rather than treating cybervictimization as the result of isolated technical failures or individual traits, the model conceptualizes it as a structured process shaped by everyday digital practices embedded in broader social contexts.

At the upstream level of the model, Social Digital Pressure (SDP) and Social Support (SS) are positioned as social-contextual antecedents. These constructs capture complementary dimensions of individuals' social environments. Social Digital Pressure reflects the extent to which digitally mediated social contexts impose expectations of constant availability, responsiveness, and visibility, while Social Support represents the availability of emotional and relational resources that can buffer stress and reduce reliance on compensatory digital behaviours. Although conceptually distinct, SDP and SS operate at the same analytical level, as both describe environmental and relational conditions rather than individual behaviours.

These social-contextual factors influence individuals' engagement with digital technologies through their impact on behavioural vulnerability patterns, most notably Smartphone Addiction (SA) and Online Gambling Behaviours. Smartphone addiction is conceptualized as a maladaptive pattern of excessive and poorly regulated device use, often driven by social expectations, emotional regulation needs, and habitual engagement. Online gambling represents a specific form of risky digital behaviour characterised by repeated financial transactions, high engagement intensity, and exposure to complex online ecosystems. In the proposed model, SA and gambling are not treated as isolated outcomes, but as behavioural mechanisms through which social conditions translate into increased exposure to digital risk.

At a more proximal level, these behavioural patterns shape individuals' lifestyle and routine digital activities, as conceptualized by the Lifestyle-Routine Activity Theory. Within the LRAT framework, exposure, proximity, target suitability, and capable guardianship jointly define the situational conditions under which cybercrime can occur. Higher levels of smartphone addiction and online gambling are expected to intensify routine digital activities, increase interaction with risky online environments, amplify visibility and attractiveness as targets, and potentially weaken effective guardianship through habituation or overconfidence. In this sense, LRAT functions as the situational mechanism linking behavioural vulnerabilities to cybervictimization risk.

Finally, Cybervictimization Risk is specified as the ultimate outcome of the model. The model allows for both direct effects of behavioural variables on cybervictimization, as well as indirect effects mediated by LRAT, reflecting the cumulative and structured nature of digital risk exposure.

Based on the theoretical framework and the empirical literature reviewed in this chapter, the following hypotheses are formulated (Table 2.1):

Table 2.1: *Model hypothesis*

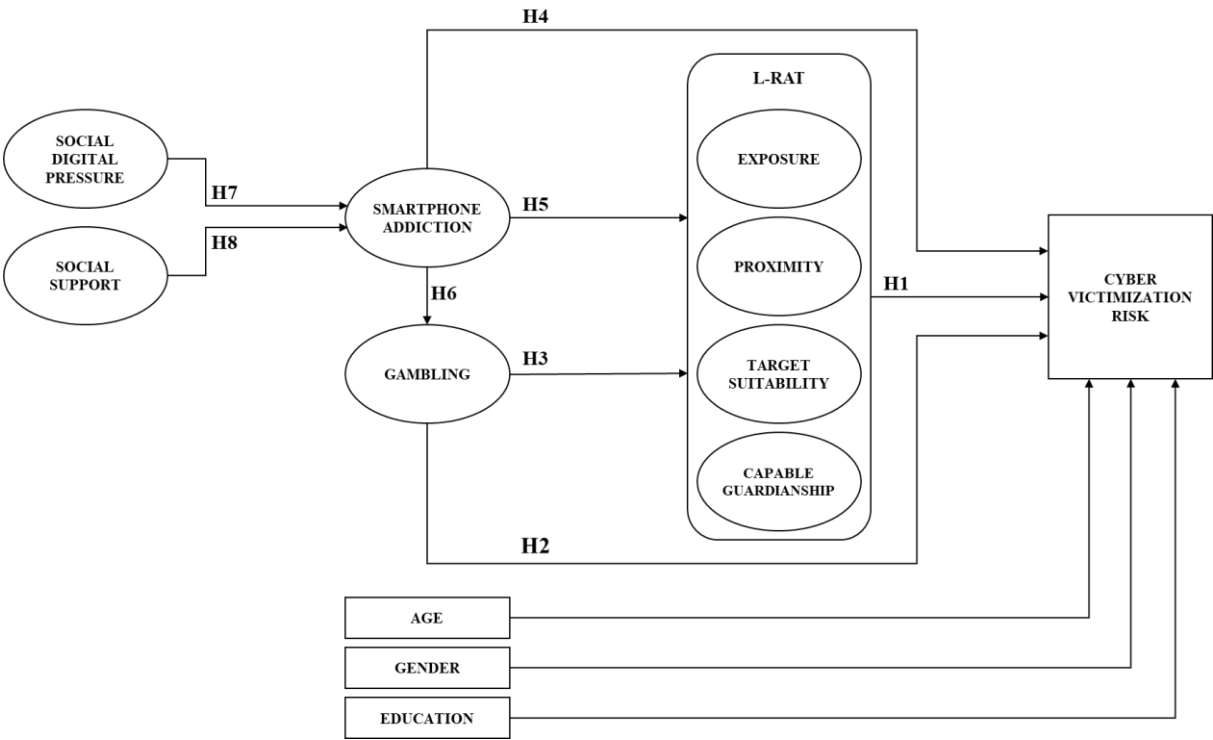
H1	Lifestyle-Routine Activity patterns are positively associated with cybervictimization risk.
H2	Online gambling behaviours are positively associated with cybervictimization risk.
H3	Online gambling behaviours indirectly increase cybervictimization risk through the L-RAT framework.
H4	Smartphone addiction is positively associated with cybervictimization risk.
H5	Smartphone addiction indirectly increases cybervictimization risk through the L-RAT framework.
H6	Smartphone addiction indirectly increases cybervictimization risk through online gambling behaviours.
H7	Higher levels of Social Digital Pressure are positively associated with Smartphone Addiction.
H8	Lower levels of perceived Social Support are negatively associated with Smartphone Addiction.

To ensure robustness and reduce potential confounding effects, the model includes a set of socio-demographic control variables commonly used in cybercrime and digital behaviour research: gender, age, and educational level. These variables are specified as direct predictors of cybervictimization and, where theoretically justified, of key behavioural constructs. Their inclusion allows the estimated relationships

between the core theoretical constructs to be interpreted as net effects, independent of basic socio-demographic differences.

Figure 2.1 displays the conceptual structural model underlying the empirical analysis, showing the hypothesised paths between antecedent social factors, behavioural constructs, the higher-order LRAT construct, and cybervictimization risk.

Figure 2.1: Structural model



2.7. Contribution of this study

2.7.1. Scientific contribution

This study contributes to cybervictimization research by advancing an integrative explanatory framework that combines criminological theory with behavioural and psychosocial perspectives. By embedding cybervictimization within Lifestyle-Routine Activity Theory and explicitly modelling the role of smartphone addiction and online gambling behaviours, the study moves beyond fragmented explanations focused on isolated risk factors. It provides empirical support for understanding cybervictimization as a socially and behaviourally embedded process shaped by everyday digital routines.

In addition, the study extends the application of LRAT to contemporary digital environments by operationalising its core components within a smartphone-centred ecosystem, contributing to the growing literature on opportunity-based explanations of cybercrime.

2.7.2. Methodological contribution

From a methodological standpoint, the study introduces two key innovations. First, it models LRAT as a second-order formative construct, consistent with its theoretical logic and allowing the joint contribution of exposure, proximity, target suitability, and guardianship to be assessed without imposing reflective assumptions. This approach responds to ongoing methodological debates in PLS-SEM regarding the appropriate specification of complex behavioural constructs.

Second, the study complements PLS-SEM estimation with machine learning techniques, enabling a comparative assessment of explanatory and predictive performance. This dual strategy strengthens the robustness of findings and aligns with recent calls for integrating confirmatory and predictive modelling in behavioural and criminological research.

2.7.3. Policy relevance

The findings of this study have direct implications for digital safety and prevention strategies. By identifying smartphone addiction, online gambling, and social digital pressure as upstream risk factors, the results suggest that effective cybercrime prevention cannot rely solely on technical safeguards. Instead, interventions should address behavioural regulation, platform design, and social expectations embedded in digital environments.

These insights are aligned with international policy frameworks, including the United Nations Agenda 2030, particularly goals related to well-being, digital inclusion, and safe technological development. The study also supports evidence-based approaches to digital security strategies at the institutional level, emphasising the importance of behavioural awareness, education, and platform responsibility in reducing cybervictimization risk.

3 Methodology

This chapter presents the methodological approach adopted to address the research objectives and empirically test the conceptual model developed in Chapter 2. In line with the study's dual emphasis on explanation and prediction, the empirical strategy combines Partial Least Squares Structural Equation Modelling (PLS-SEM) with a complementary Machine Learning approach based on Random Forest classification. The two methods are used with distinct but convergent purposes: PLS-SEM supports theory-driven estimation and interpretation of relationships among latent constructs, whereas Machine Learning is employed to evaluate predictive performance and identify the most influential predictors in a risk-classification setting.

The chapter is structured as follows. First, regression-type methods are briefly introduced together with their main limitations in addressing complex behavioural models involving latent constructs, mediation mechanisms, and measurement error. The chapter then outlines the Structural Equation Modelling framework, clarifying the distinction between covariance-based SEM and variance-based PLS-SEM, and justifying the adoption of PLS-SEM given the objectives and characteristics of the present research design. Particular attention is given to the components of a PLS-SEM model and to the evaluation procedures required to ensure measurement quality and structural validity, including the assessment of reflective and formative measurement models, higher-order constructs, and mediation effects.

Finally, the chapter introduces the Machine Learning framework and motivates its integration alongside PLS-SEM. In this study, cybervictimization risk is modelled as a continuous index in the explanatory PLS-SEM analysis, consistent with the objective of capturing graded differences in vulnerability. By contrast, for predictive modelling, the outcome is operationalized as a binary target (victimized vs non-victimized) to reflect a classification logic that is more directly aligned with risk identification and predictive accuracy. This methodological choice allows the study to jointly address explanatory questions (why and through which mechanisms cybervictimization risk increases) and predictive questions (how accurately risk profiles can be classified and which predictors contribute most to prediction).

3.1. Regression types methods and their limits

Regression-type methods refer to a broad family of first-generation statistical techniques designed to estimate the relationship between a dependent variable and one or more independent variables by minimizing the discrepancy between observed and predicted values, typically through Ordinary Least Squares (OLS) estimation (Kutner, 2005; Haenlein & Kaplan, 2004). Formally, regression analysis aims to model the conditional expectation of a response variable as a function of one or more predictors, under a set of assumptions regarding linearity, independence, and measurement properties of the variables involved (Kutner, 2005).

In applied research, the expression regression-type methods commonly includes all traditional multivariate techniques based on OLS or its extensions, such as multiple linear regression, logistic regression, and analysis of variance (ANOVA), which have historically represented the core empirical toolkit in the social and behavioural sciences (Haenlein & Kaplan, 2004).

3.1.1. Examples of regression-type methods

Several well-established statistical techniques fall within the category of regression-type methods.

Multiple linear regression is a technique used to estimate the linear relationship between a continuous dependent variable and multiple independent variables, assuming additive and linear effects of predictors on the outcome (Kutner, 2005). It is widely applied when the objective is to assess the relative contribution of several predictors to a single observed outcome.

Logistic regression extends the regression framework to situations in which the dependent variable is categorical, most commonly binary. Instead of modelling the outcome directly, logistic regression estimates the probability of occurrence of an event through a logistic function, linking predictors to the log-odds of the response variable (Hosmer et al., 2013).

Analysis of Variance (ANOVA) is a statistical method designed to test whether the means of a continuous dependent variable differ significantly across two or more categorical groups. ANOVA decomposes total variance into between-group and within-group components, allowing the assessment of group-level effects under specific distributional assumptions (Field, 2024).

Multiple ANOVA (or factorial ANOVA) extends the basic ANOVA framework by allowing the simultaneous analysis of multiple categorical independent variables and their interaction effects on a continuous dependent variable. This approach is useful for examining how combinations of factors jointly influence an outcome, although it remains limited to relatively simple model structures (Field, 2024).

3.1.2. Limitations of regression-type methods

Despite their widespread adoption and ease of implementation, regression-type methods present several structural limitations that restrict their suitability for complex behavioural research designs.

The first limitation concerns model complexity. Regression-based techniques are inherently designed to estimate relatively simple model structures involving a single dependent variable and a set of predictors. While indirect effects or mediation structures can be approximated through sequential regressions, this piecemeal approach reduces statistical efficiency and increases the risk of biased parameter estimates. A causal chain such as “A leads to B leads to C” can technically be modelled using multiple regressions, but cannot be estimated simultaneously, limiting coherence and interpretability (Haenlein & Kaplan, 2004).

A second limitation relates to the nature of the variables included in the model. Regression-type methods require variables to be directly observable, such as age, income, or frequency measures. However, many key constructs in behavioural and social sciences, such as addiction, perceived pressure, or psychological vulnerability, are latent by nature and cannot be directly observed. In traditional regression frameworks, these theoretical constructs can only be incorporated after separate validation procedures, for example through confirmatory factor analysis, which fragments the analytical process (Bagozzi & Phillips, 1982).

The third major limitation concerns measurement error. First-generation techniques implicitly assume that all variables are measured without error, an assumption that is rarely met in real-world behavioural data. Self-reported measures, subjective perceptions, and survey-based indicators are inevitably affected by random and systematic measurement error. Regression models do not explicitly account for this source of bias, potentially leading to attenuated or distorted estimates. As emphasized in the methodological literature, these techniques are optimal only under idealized conditions that do not reflect the complexity of social and behavioural phenomena (Haenlein & Kaplan, 2004).

For these reasons regression-type methods are inadequate for addressing the multidimensional, latent, and interrelated constructs examined in the present study, motivating the adoption of structural equation modelling as a more appropriate analytical framework.

3.2. Structural Equation Modelling

3.2.1. SEM: foundations and rationale for the adoption of PLS-SEM

Structural Equation Modelling (SEM) is a family of multivariate techniques that enables the simultaneous estimation of multiple interdependent relationships among

observed and latent variables. It is widely employed across psychology, behavioural sciences, sociology, and business research because it supports theory-driven modelling of complex causal structures, including mediation and indirect effects, while explicitly incorporating measurement error (Bagozzi & Phillips, 1982; Hair et al., 2019). Compared with first-generation methods such as multiple regression or ANOVA, SEM is particularly advantageous when the research framework involves latent constructs that cannot be directly observed, when indicators are imperfect proxies of theoretical concepts, or when the hypothesized relationships form a system of equations rather than a single dependent outcome (Haenlein & Kaplan, 2004).

A defining characteristic of SEM is the separation between (i) the measurement model, which specifies how latent constructs are operationalized through observed indicators, and (ii) the structural model, which encodes the hypothesized relationships among the constructs. This dual architecture strengthens the link between conceptual theory and empirical estimation by allowing researchers to test structural hypotheses while explicitly accounting for measurement quality, an aspect that traditional regression frameworks typically treat implicitly or ignore (Bagozzi & Phillips, 1982).

Within SEM, two main methodological traditions are typically distinguished: covariance-based SEM (CB-SEM) and partial least squares SEM (PLS-SEM). Although both share the same underlying logic of linking measurement and structural components, they diverge in estimation goals, assumptions, and evaluation criteria (Hair et al., 2019; Rigdon et al., 2017).

CB-SEM is primarily oriented toward theory testing and confirmation. Its objective is to assess the extent to which a theoretically specified model can reproduce the observed covariance matrix, typically through maximum likelihood (ML) or related estimators (Joreskog, 1971). Consequently, CB-SEM places strong emphasis on global model fit, using fit indices and residual-based diagnostics to evaluate whether the proposed theory is consistent with the empirical covariance structure (Hair et al., 2019). This confirmatory orientation implies relatively strict requirements: CB-SEM commonly assumes multivariate normality (or the use of robust alternatives when normality is violated), demands adequate sample size for stable covariance estimation and reliable fit testing, and is sensitive to misspecification, especially when models are complex or include many constructs and indicators. In practice, complex CB-SEM models may encounter estimation difficulties (e.g.: non-convergence, inadmissible solutions such as negative variance estimates, or identification problems), particularly when data quality, distributional properties, or sample size are not optimal (Hair et al., 2019). For these reasons, CB-SEM is most appropriate when theoretical structures are well established, measurement models are mature, and the research goal is rigorous confirmation of a prespecified model rather than maximizing predictive performance.

In contrast, PLS-SEM is typically framed as a causal-predictive approach that prioritizes explained variance and prediction of endogenous constructs over exact covariance reproduction (Hair et al., 2019). PLS-SEM uses an iterative algorithm based on a sequence of ordinary least squares regressions to estimate latent variable scores and path relationships, aiming to maximize the R^2 values of endogenous constructs. As a result, PLS-SEM is often preferred when the research objective emphasizes prediction, variance explanation, and model development in emerging domains where theory is still evolving (Hair et al., 2019; Henseler, 2018). Compared to CB-SEM, PLS-SEM is generally more tolerant of non-normal data and can be advantageous in settings with complex models, many constructs, or formative measurement specifications (Tenenhaus et al., 2005). Moreover, empirical comparisons suggest that PLS-SEM can be particularly practical under certain sample-size conditions and model structures, although the choice between ML-based and PLS-based estimators should remain aligned with the study's aim rather than treated as universally superior (Barroso et al., 2010; Rigdon et al., 2017).

It is important, however, to avoid portraying PLS-SEM as a “one-size-fits-all” solution. Methodological debates emphasize that PLS-SEM is not a “silver bullet,” and that researchers should justify its adoption based on a transparent alignment between research goals, measurement characteristics, and evaluation strategy (Hair et al., 2011; Rigdon et al., 2017). Contemporary guidance stresses the need to apply appropriate quality criteria for measurement models (reflective vs. formative), report structural results with adequate inferential procedures (e.g.: bootstrapping), and interpret predictive performance with metrics consistent with the causal-predictive orientation (Hair et al., 2019, 2021).

The adoption of PLS-SEM in the present study is justified by three main considerations. First, the research framework integrates multiple latent constructs, such as smartphone addiction, social digital pressure, social support, and lifestyle/routine activity components, that are not directly observable and are operationalized through multiple indicators. PLS-SEM enables the explicit modelling of these constructs while accounting for measurement error, and it offers flexible handling of measurement specifications, including the possibility of formative relationships where theoretically appropriate (Tenenhaus et al., 2005; Hair et al., 2019). Second, the study specifies a complex structural model characterized by multiple interrelated paths and indirect effects. In such contexts, PLS-SEM can provide stable estimation and efficient variance explanation, aligning with the objective of understanding how psychosocial and behavioural factors jointly relate to cybervictimization risk (Hair et al., 2019; Rigdon et al., 2017). Third, the exploratory and predictive orientation of the research, focused on assessing explanatory and predictive relationships within a multifaceted framework, fits the typical strengths of PLS-SEM as highlighted in state-of-the-art methodological discussions and reporting guidelines (Henseler, 2018; Hair et al., 2019, 2021).

3.2.2. Components of a PLS-SEM

Partial Least Squares Structural Equation Modelling (PLS-SEM) is articulated through a set of interdependent elements that translate theoretical concepts into an empirical model and enable the statistical assessment of hypotheses. At its core, a PLS-SEM specification comprises indicators, latent constructs, and relationships among these constructs. These elements are organized into two connected submodels: the measurement model (outer model) and the structural model (inner model) (Hair et al., 2019, 2021). This explicit separation is a defining feature of SEM approaches, as it allows researchers to evaluate measurement quality independently from the substantive theoretical relationships, thereby improving interpretability and construct validity relative to regression-only frameworks (Bagozzi & Phillips, 1982).

3.2.2.1. Indicators, constructs, and relationships

PLS-SEM distinguishes between observable indicators and latent constructs. Indicators are directly measured variables, commonly survey items, behavioural frequencies, or digital trace measures, that provide empirical information about an underlying theoretical domain. Constructs, in turn, represent abstract entities (e.g.: attitudes, dispositions, tendencies, or psychosocial states) that cannot be observed directly and must therefore be inferred from patterns among multiple indicators (Bagozzi & Phillips, 1982; Hair et al., 2019).

Relationships in PLS-SEM operate at two conceptual levels. At the measurement level, relationships define how indicators connect to their construct, thereby specifying the operationalization of the theoretical concept (outer model). At the structural level, relationships are specified as directional paths among constructs, reflecting theoretically grounded hypotheses about causal or predictive effects (inner model) (Hair et al., 2019, 2021).

This dual structure enables researchers to disentangle problems of measurement (e.g.: weak indicators or mis specified measurement direction) from substantive theoretical claims (e.g.: the magnitude and significance of paths), strengthening both inference and theoretical interpretation.

3.2.2.2. Measurement model (outer model)

The measurement model specifies how each construct is measured by its indicators. In PLS-SEM, this specification can follow either a reflective or a formative logic, depending on the theorized direction of causality between the construct and its measures (Jarvis et al., 2003; Hair et al., 2019).

In reflective measurement models, the construct is theorized to drive the observed indicators. Changes in the latent construct are expected to manifest as systematic changes across its indicators; accordingly, indicators typically show shared variance and are often treated as substitutable manifestations of the same underlying

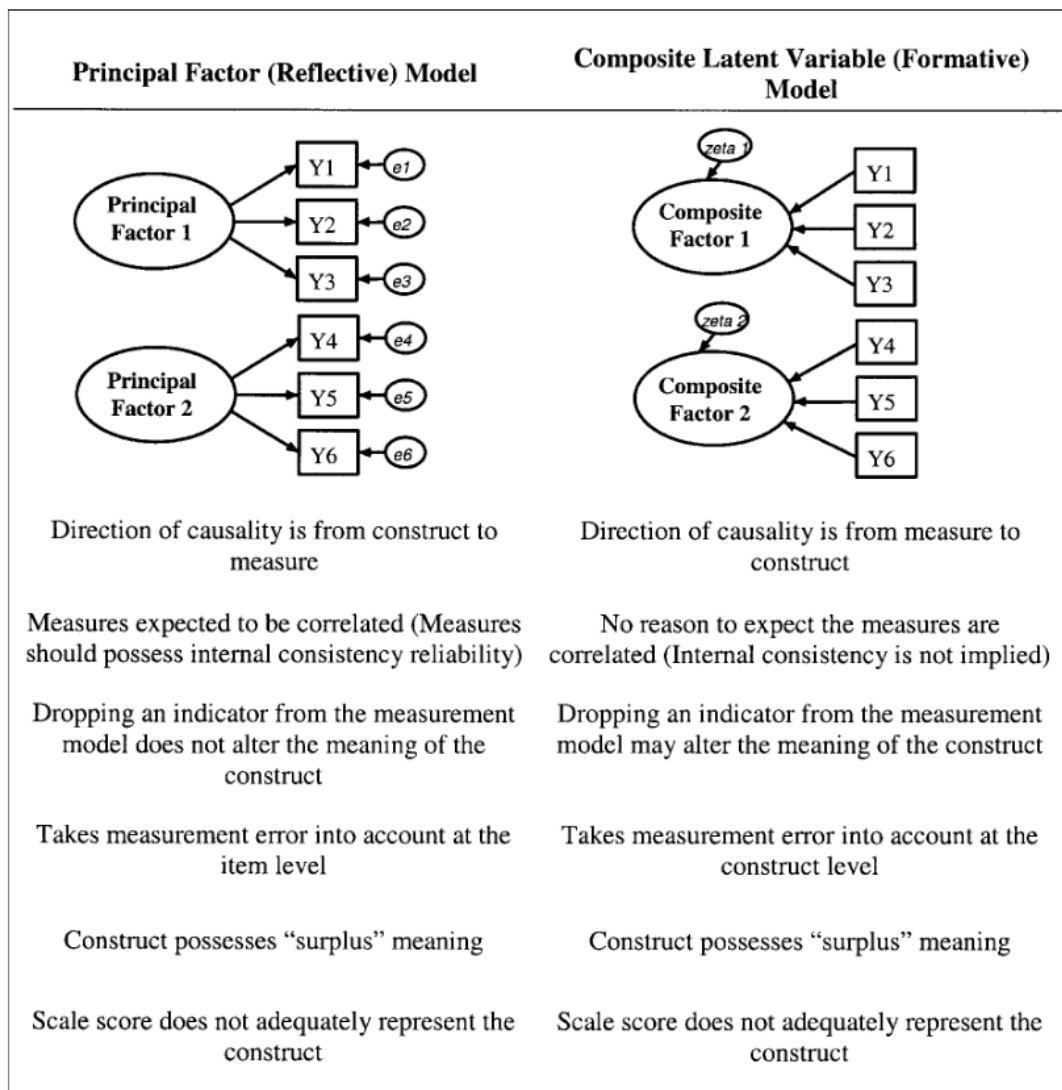
concept. Because the construct is the common cause, dropping one indicator generally does not redefine the construct's conceptual meaning, provided that the remaining indicators still capture the domain adequately (Jarvis et al., 2003). Consistent with this logic, reflective measurement evaluation focuses on internal consistency reliability, convergent validity, and discriminant validity, using criteria such as composite reliability, average variance extracted (AVE), and appropriate discriminant validity assessments. Reflective specifications are typically appropriate when indicators are conceived as parallel expressions of an underlying psychological or attitudinal attribute, for instance, when the theory suggests that a shift in the latent trait should be mirrored across multiple items capturing the same disposition (Hair et al., 2019; Henseler et al., 2015; Henseler, 2018).

In formative measurement models, the causal logic is reversed: indicators jointly define or “compose” the construct. Each indicator represents a distinct dimension that contributes to the construct's meaning, and the indicators are not required to correlate. In this case, removing an indicator may alter the conceptual coverage of the construct, because the construct is defined by the set of its indicators (Jarvis et al., 2003). For formative models, traditional internal consistency metrics (e.g.: Cronbach's alpha) are not appropriate, since the indicators need not share variance by design. Instead, formative measurement evaluation typically emphasizes (i) the relevance and significance of indicator weights, (ii) collinearity diagnostics among indicators (e.g.: VIF), and (iii) convergent validity checks such as redundancy analysis when applicable (Hair et al., 2019, 2021). Formative specifications are particularly suitable for constructs that are deliberately assembled from non-interchangeable components, especially when theory defines the construct as an aggregation of distinct dimensions rather than as a single underlying trait; this logic is often relevant for framework-driven composites, such as those derived from lifestyle/routine activity components (Jarvis et al., 2003; Hair et al., 2021).

Correctly specifying reflective versus formative measurement is critical. Misspecification, especially treating formative measures as reflective (or vice versa), can bias parameters, distort validity assessment, and ultimately lead to incorrect substantive conclusions (Jarvis et al., 2003; Hair et al., 2019).

Figure 3.1 summarizes the key conceptual and diagnostic differences between reflective and formative measurement models, highlighting how the direction of causality and indicator properties change across the two specifications.

Figure 3.1: Summary of differences between types of measurement models (Source: Jarvis et al., 2003)



3.2.2.3. Structural model (inner model)

The structural model represents the system of hypothesized relationships among latent constructs and constitutes the theoretical core of the analysis. Structural paths encode direct effects between constructs and are specified based on theoretical reasoning and prior empirical evidence. In PLS-SEM, the structural model is estimated with a variance-based objective: rather than reproducing the observed covariance matrix, the approach prioritizes maximizing the explained variance of designated dependent constructs (Tenenhaus et al., 2005; Hair et al., 2019).

A key distinction in the structural model is between exogenous and endogenous constructs. Exogenous constructs function as predictors only: they have outgoing paths but no incoming structural paths. Endogenous constructs are explained by the model: they have one or more incoming paths, and their variance is partially

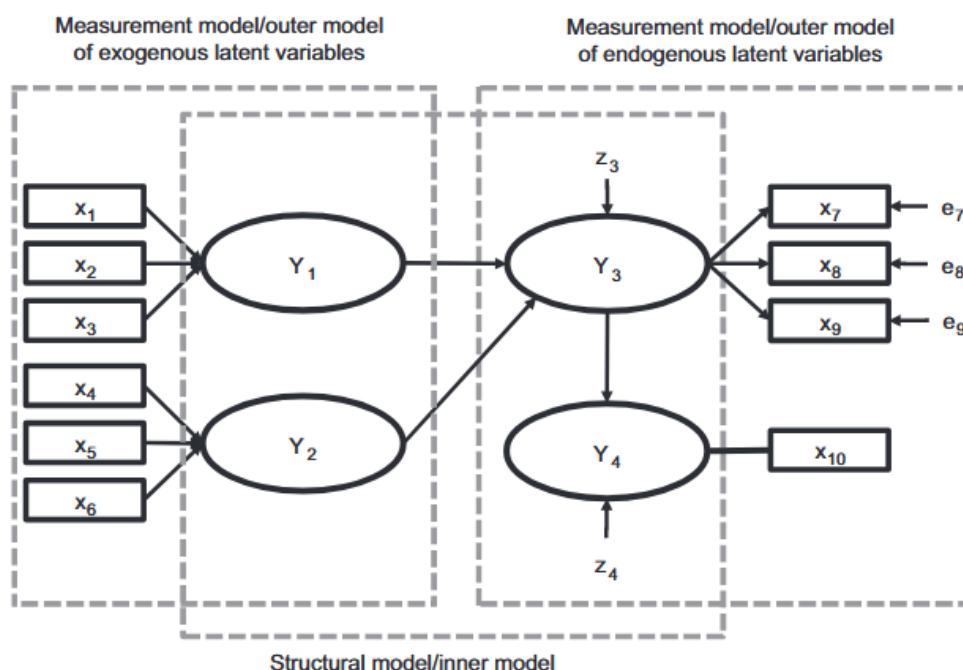
accounted for by antecedent constructs. This distinction is substantively important because model evaluation in PLS-SEM is largely centred on the predictive/explanatory performance for endogenous constructs (Hair et al., 2019, 2021).

Structural model assessment in PLS-SEM focuses on evaluating whether the hypothesized relationships among constructs are empirically supported and whether the model achieves an adequate level of explanatory power. In this context, attention is primarily directed toward the estimated structural paths, which capture the direction and relative strength of the relationships among constructs, as well as toward the ability of the model to explain variance in the endogenous constructs. This explanatory capability is commonly summarized through the coefficient of determination (R^2), which indicates the proportion of variance in an endogenous construct accounted for by its predictors in the model (Hair et al., 2019, 2021). In addition, structural model evaluation requires ensuring that predictor constructs are not affected by excessive collinearity, as high intercorrelations among antecedent variables may compromise the stability and interpretability of path estimates.

Beyond R^2 , many PLS-SEM workflows also consider additional predictive-oriented metrics (e.g.: out-of-sample predictive assessment) depending on research goals; however, the foundational evaluation of the inner model is anchored in path estimates and variance explanation (Hair et al., 2019; Henseler, 2018).

Figure 3.2 provides an integrative representation of the PLS-SEM architecture discussed above, explicitly visualizing how indicators relate to constructs in the outer model and how constructs are linked through hypothesized paths in the inner model, thereby offering a compact synthesis of indicators, constructs, and relationships within a single modelling framework.

Figure 3.2: General structure of a PLS-SEM model (Source: J. F. Hair et al., 2021)



3.2.2.4. Bootstrapping in PLS-SEM

Because PLS-SEM does not depend on strong distributional assumptions, classical parametric inference based on normal-theory standard errors is often inappropriate. Consequently, hypothesis testing and uncertainty quantification in PLS-SEM are commonly performed using bootstrapping, a nonparametric resampling procedure that approximates the sampling distribution of model estimates. In bootstrapping, a large number of resamples are drawn with replacement from the original dataset; the PLS-SEM model is re-estimated in each resample, generating an empirical distribution for each parameter (e.g.: path coefficients, outer loadings, and formative weights) (Efron & Tibshirani, 1994; Hair et al., 2019, 2021).

This empirical distribution supports several inferential outputs. First, it enables the computation of standard errors, t-values, and p-values for parameters, which are then used to evaluate the statistical significance of hypothesized effects in the structural model and the relevance of indicators in the measurement model (Hair et al., 2019, 2021). Second, and increasingly emphasized in methodological recommendations, bootstrapping provides confidence intervals (CIs) as a direct basis for statistical inference. In this setting, a parameter is typically considered statistically significant at a given level if the corresponding CI does not include zero (Wood, 2005; Streukens & Leroi-Werelds, 2016).

Confidence interval choice matters. Compared with basic percentile intervals, bias-corrected percentile confidence intervals adjust for potential bias in the bootstrap distribution and often provide more accurate coverage, particularly in finite samples or when bootstrap distributions are asymmetric (Streukens & Leroi-Werelds, 2016; Wood, 2005). Accordingly, PLS-SEM reporting guidance commonly recommends presenting bootstrap-based CIs, preferably bias-corrected, alongside or in place of sole reliance on t-tests, because CIs offer a more informative description of estimation uncertainty (Hair et al., 2019, 2021).

Practical implementation also requires methodological care. Researchers must specify an adequate number of bootstrap resamples to stabilize standard errors and CI estimates, and they should apply consistent settings (e.g.: two-tailed vs one-tailed testing aligned with the directional nature of hypotheses). Moreover, when interpreting bootstrap results, it is important to examine not only statistical significance but also effect magnitude and the plausibility of estimated signs, particularly in complex models with multiple paths and potential collinearity (Hair et al., 2019, 2021; Streukens & Leroi-Werelds, 2016).

In sum, bootstrapping is not an optional add-on but a core inferential mechanism in PLS-SEM. It enables robust significance assessment for both measurement and structural parameters and supports transparent reporting through confidence intervals, thereby strengthening the credibility of the empirical results presented in this study.

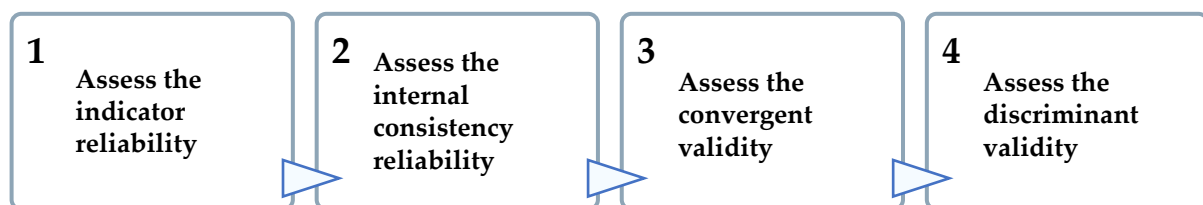
3.2.3. Evaluation of Reflective Measurement Models

The evaluation of reflective measurement models constitutes a critical stage in PLS-SEM analysis, as it establishes whether latent constructs are measured in a reliable and valid manner prior to examining the structural relationships among them. Methodological guidelines consistently emphasize that inadequately specified or poorly performing measurement models can bias parameter estimates and undermine substantive interpretations of the structural model (Hair et al., 2019, 2021). For this reason, reflective measurement model assessment is conducted before any evaluation of the inner model.

In line with established recommendations, the assessment of reflective measurement models follows a sequential procedure encompassing indicator reliability, internal consistency reliability, convergent validity, and discriminant validity (Hair et al., 2021). This ordered process ensures that indicators appropriately represent their intended constructs and that constructs are empirically distinct from one another, thereby providing a sound foundation for subsequent hypothesis testing.

Figure 3.3 illustrates this stepwise assessment procedure and summarizes the main stages involved in evaluating reflective measurement models in PLS-SEM.

Figure 3.3: Reflective measurement model assessment procedure (adapted from J. F. Hair et al., 2021)



3.2.3.1. Indicator reliability

Indicator reliability refers to the degree to which an observed indicator is explained by its associated latent construct. In reflective measurement models, indicators are assumed to be manifestations of the underlying construct; therefore, a substantial proportion of each indicator's variance should be attributable to the latent variable. Empirically, indicator reliability is assessed through outer loadings, which represent the standardized bivariate relationship between an indicator and its construct.

The square of an outer loading reflects the proportion of indicator variance explained by the construct. As a commonly accepted guideline, an outer loading of 0.70 or higher implies that at least 50% of the indicator's variance is captured by the latent construct, which is considered adequate for reliable measurement (Hair et al., 2021). Indicators with loadings above this threshold contribute substantially to construct measurement and are typically retained without further scrutiny.

Indicators with outer loadings between 0.40 and 0.70 require a more detailed examination. According to methodological recommendations, such indicators may be retained if their removal does not lead to a meaningful improvement in internal consistency reliability (2nd step) or convergent validity (3rd step). This consideration is particularly relevant in exploratory research or when theoretical content validity would be compromised by indicator deletion. By contrast, indicators with very low outer loadings generally explain more error variance than construct variance and are therefore candidates for removal, as they can weaken the overall quality of the measurement model (Hair et al., 2012, 2021).

3.2.3.2. Internal consistency reliability

Internal consistency reliability assesses whether the set of indicators associated with a construct consistently measures the same underlying concept. In PLS-SEM, reliability assessment does not rely exclusively on traditional assumptions of equal indicator loadings (tau-equivalence). Instead, it adopts reliability measures that explicitly account for differences in indicator contributions (Hair et al., 2019, 2021).

The primary reliability measure is composite reliability (ρ_C), which incorporates the actual outer loadings and provides a more accurate estimate than Cronbach's alpha. However, current standards also recommend reporting ρ_A , a Dijkstra-Henseler reliability coefficient that often represents a more precise "middle ground" between the conservative Alpha and the potentially liberal ρ_C . Values of 0.70 or higher for ρ_C and ρ_A are generally considered acceptable. While Cronbach's alpha is reported for descriptive completeness, it is regarded as a conservative lower-bound estimate because it assumes equal indicator loadings. Consequently, reliability should be validated across all three indices, ensuring they stay below 0.95 to avoid item redundancy (Hair et al., 2019, 2021).

3.2.3.3. Convergent validity of reflective constructs

Convergent validity evaluates the extent to which a latent construct explains the variance of its indicators relative to measurement error. In reflective measurement models, convergent validity is assessed using the Average Variance Extracted (AVE), which represents the mean proportion of variance that a construct captures from its indicators (Fornell & Larcker, 1981).

An AVE value of 0.50 or higher indicates that, on average, the construct explains more than half of the variance of its indicators, thereby providing empirical support for convergent validity. AVE values below this threshold imply that measurement error dominates the explained variance, raising concerns about the adequacy of the construct's operationalization (Hair et al., 2019, 2021). In such cases, researchers are encouraged to reconsider indicator selection or construct specification.

Establishing convergent validity is essential for meaningful interpretation of structural relationships, as constructs that fail to adequately capture their indicators can attenuate path coefficients and obscure substantive effects in the structural model.

3.2.3.4. Discriminant validity

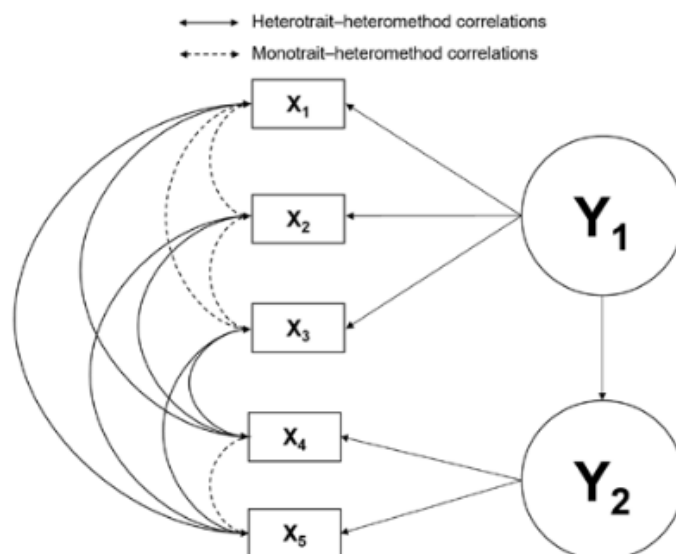
Discriminant validity assesses whether a construct is empirically distinct from other constructs in the model. Adequate discriminant validity ensures that each construct captures a unique theoretical concept and that observed relationships among constructs are not inflated by measurement overlap (Hair et al., 2019, 2021).

Traditionally, discriminant validity has been evaluated using the Fornell-Larcker criterion, which compares the square root of a construct's AVE with its correlations with other constructs. Discriminant validity is supported when the square root of the AVE exceeds the construct's correlations with all other constructs, indicating that the construct shares more variance with its indicators than with other latent variables (Fornell & Larcker, 1981).

Subsequent research has demonstrated that the Fornell-Larcker criterion may lack sensitivity in detecting discriminant validity problems, particularly in models with closely related constructs. To address this limitation, the Heterotrait-Monotrait ratio of correlations (HTMT) has been proposed as a more stringent and reliable criterion (Henseler et al., 2015). HTMT evaluates discriminant validity by comparing the average correlations across constructs with the average correlations within the same construct. Discriminant validity is established when HTMT values fall below recommended thresholds, typically 0.85 for conceptually distinct constructs and up to 0.90 for constructs that are theoretically related (Hair et al., 2019, 2021).

Figure 3.4 presents a graphical representation of the HTMT criterion and highlights the distinction between heterotrait-heteromethod and monotrait-heteromethod correlations.

Figure 3.4: Discriminant validity assessment using the HTMT
(Source: J. F. Hair et al. 2021)



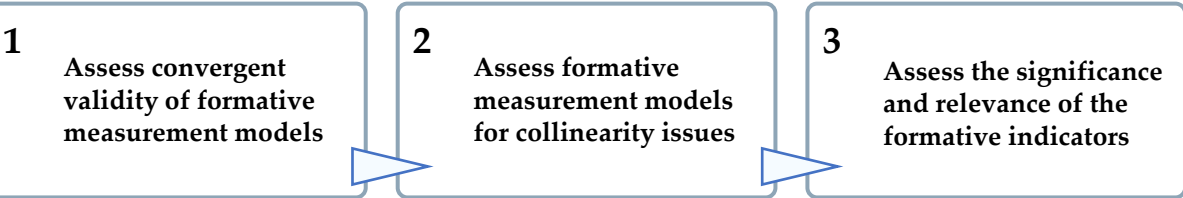
Given their complementary strengths, current methodological guidelines recommend applying both the Fornell-Larcker criterion and HTMT when assessing discriminant validity in reflective measurement models, especially in applied and complex research settings (Hair et al., 2019, 2021).

3.2.4. Evaluation of Formative Measurement Models

The evaluation of formative measurement models differs substantially from that of reflective models, as formative indicators are not assumed to be manifestations of an underlying latent construct but instead jointly define and compose the construct itself. As a consequence, traditional reliability and validity criteria commonly applied to reflective constructs, such as internal consistency reliability or average variance extracted, are not meaningful in the formative context.

Rather than assessing internal consistency, the evaluation of formative measurement models follows a distinct and theoretically grounded logic that focuses on three key aspects: convergent validity, indicator collinearity, and the relevance and significance of indicator weights. Together, these criteria ensure that formative constructs are conceptually well specified, empirically meaningful, and statistically stable (Hair et al., 2019, 2021).

Figure 3.5: Formative measurement model assessment procedure (adapted from J. F. Hair et al., 2021)



3.2.4.1. Convergent validity of formative constructs

Convergent validity represents the first step in the evaluation of formative measurement models. Unlike reflective constructs, convergent validity in formative models is assessed through redundancy analysis, which examines the extent to which a formative construct correlates with an alternative measure of the same concept specified as reflective.

In practical terms, redundancy analysis involves linking the formative construct to a global single-item or short reflective measure that captures the same theoretical domain. The strength of the relationship between the formative construct and its reflective proxy is then evaluated. Convergent validity is considered established when the path coefficient between the formative construct and the reflective measure exceeds the threshold of 0.708, indicating that the formative construct explains more than 50% of the variance in the reflective proxy (Hair et al., 2019, 2021).

The availability of an appropriate reflective proxy is a critical prerequisite for redundancy analysis. When such a proxy is not available, convergent validity cannot be empirically assessed, and the evaluation of the formative construct must rely primarily on strong theoretical justification and content validity considerations.

3.2.4.2. Indicator collinearity

The second criterion for evaluating formative measurement models concerns collinearity among indicators. Because formative indicators jointly define a construct, excessive collinearity can distort the estimation of indicator weights, inflate standard errors, and compromise the interpretability and stability of the measurement model.

Collinearity is assessed using the Variance Inflation Factor (VIF), which quantifies the extent to which an indicator is linearly related to the other indicators forming the same construct. As a general guideline, VIF values should remain below the threshold of 5, while more conservative recommendations suggest values below 3 to ensure the absence of critical collinearity issues (Diamantopoulos & Winklhofer, 2001; Hair et al., 2019, 2021).

Indicators exhibiting excessive VIF values may signal redundancy or conceptual overlap and therefore require careful examination. However, decisions regarding indicator removal must always be theoretically justified, as eliminating formative indicators may alter the conceptual meaning and content validity of the construct.

3.2.4.3. Relevance and significance of indicator weights

The final step in the evaluation of formative measurement models concerns the relevance and statistical significance of indicator weights. Indicator weights reflect the relative contribution of each indicator to the formation of the construct and provide insight into the importance of individual dimensions within the formative construct.

The statistical significance of indicator weights is assessed using bootstrapping procedures, which generate empirical confidence intervals without relying on distributional assumptions. An indicator weight is considered statistically significant when its confidence interval does not include zero. Significant weights indicate that the indicator contributes uniquely to the construct, whereas non-significant weights suggest a weaker empirical contribution (Hair et al., 2019, 2021).

Importantly, non-significant indicator weights do not automatically justify indicator removal. In formative measurement models, indicators may still be theoretically essential even when their statistical contribution is limited. In such cases, indicator relevance can be further evaluated by examining indicator loadings, which capture the absolute contribution of each indicator to the construct. Indicators characterized by non-significant weights but substantial loadings may be retained in order to preserve the conceptual completeness and theoretical integrity of the formative construct.

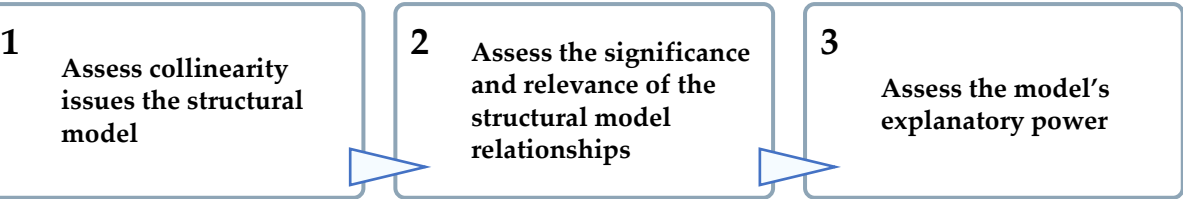
3.2.5. Evaluation of the Structural Model

Once the quality of the measurement models has been established, the evaluation of the structural model represents the subsequent step in the PLS-SEM analysis. The assessment of the structural model focuses on the relationships among latent constructs and aims to verify whether the hypothesized paths are statistically significant, theoretically meaningful, and capable of explaining a substantial proportion of variance in the endogenous constructs.

Unlike covariance-based SEM, PLS-SEM does not provide global goodness-of-fit indices, as its objective is not to reproduce the observed covariance matrix but to maximize the explained variance of the endogenous constructs. Consequently, the evaluation of the structural model relies on predictive-oriented criteria rather than overall model fit measures (Tenenhaus et al., 2005; Hair et al., 2019, 2021).

In line with established methodological guidelines, the evaluation of the structural model follows a sequential procedure that includes the assessment of collinearity among predictor constructs, the significance of path coefficients, and measures of explanatory power, such as the coefficient of determination (R^2) and effect size (f^2).

Figure 3.6: Structural model assessment procedure (adapted from J. F. Hair et al., 2021)



3.2.5.1. Collinearity

The presence of collinearity among predictor constructs represents a critical issue in the evaluation of the structural model, as high levels of collinearity can inflate standard errors, reduce statistical power, and lead to unstable and unreliable path coefficient estimates. Although collinearity is particularly emphasized in formative measurement models, it can also adversely affect the estimation and interpretation of structural relationships.

In PLS-SEM, collinearity in the structural model is assessed using the Variance Inflation Factor (VIF), computed on latent variable scores for each set of predictor constructs involved in a structural path. As a general guideline, VIF values should remain below 5 to avoid severe collinearity problems, while more conservative recommendations suggest maintaining values below 3 to ensure robust estimation (Hair et al., 2019, 2021).

When VIF values exceed acceptable thresholds, researchers should consider revising the model specification or addressing redundancy among predictor constructs, provided that such modifications are theoretically justified.

3.2.5.2. Significance of path coefficients

The second step in structural model evaluation involves assessing the statistical significance and magnitude of path coefficients, which represent the strength and direction of the hypothesized relationships among latent constructs. In PLS-SEM, statistical inference is conducted using bootstrapping procedures, as the method does not rely on distributional assumptions (Hair et al., 2019, 2021).

Bootstrapping generates empirical sampling distributions for each path coefficient, allowing the estimation of standard errors, t-values, and confidence intervals. Path coefficients are considered statistically significant when their associated t-values exceed critical thresholds (e.g.: 1.96 for a two-tailed test at the 5% significance level) or, equivalently, when bias-corrected confidence intervals do not include zero.

Path coefficients typically range between -1 and +1. Values outside this interval may indicate estimation problems, often related to excessive collinearity or model misspecification, and should therefore be carefully examined before interpretation.

3.2.5.3. Measures of explanatory power: coefficient of determination (R^2)

After evaluating individual path coefficients, the explanatory power of the structural model is assessed by examining the coefficient of determination (R^2) for each endogenous construct. In PLS-SEM, the coefficient of determination represents the proportion of variance in an endogenous latent construct that is explained by its predictor constructs. Unlike covariance-based regression models, R^2 in PLS-SEM is derived from latent variable scores and structural model estimates rather than from observed variable variance decomposition.

Formally, R^2 can be expressed as a function of the structural path coefficients and the correlations between latent constructs:

$$R^2 = \sum_j |\hat{\beta}_j \text{cor}(\eta_i, \xi_j)|$$

Equation 1: Coefficient of determination (R^2)

In this formulation, $\hat{\beta}_j$ denotes the structural path coefficient linking the exogenous construct ξ_j to the endogenous construct η_i , while $\text{cor}(\eta_i, \xi_j)$ represents the correlation between their latent variable scores. Each term in the summation captures the contribution of a predictor construct to the explained variance of the endogenous construct, and their aggregation provides a measure of overall explanatory power.

Because R^2 is not directly estimated but computed as a function of multiple model parameters, its statistical uncertainty cannot be derived analytically. For this reason,

confidence intervals for R^2 are obtained using bootstrapping procedures. Specifically, R^2 is computed separately for each of the bootstrap samples based on the corresponding path coefficients and latent variable correlations. The resulting empirical distribution of R^2 values is then used to construct a bias-corrected percentile bootstrap confidence interval, providing an assessment of the stability and robustness of the explained variance estimate. This bootstrap-based procedure for constructing confidence intervals of R^2 in PLS-SEM follows the approach originally formalized by Tenenhaus et al. (2005) and subsequently adopted in applied guidelines (Hair et al., 2019).

R^2 values range from 0 to 1, with higher values indicating greater explanatory power. As general guidelines, values of approximately 0.75, 0.50, and 0.25 may be interpreted as substantial, moderate, and weak, respectively, although these thresholds should always be considered in relation to the research context and disciplinary norms.

Despite its usefulness, R^2 presents important limitations. The coefficient can increase mechanically with the inclusion of additional predictor constructs, even when these predictors have limited theoretical relevance. Moreover, high R^2 values may result from overfitting, particularly in complex models. For these reasons, R^2 should be interpreted in conjunction with additional measures of explanatory power.

3.2.5.4. Measures of explanatory power: effect size (f^2)

To complement the interpretation of R^2 , PLS-SEM employs the effect size (f^2) to assess the relative contribution of individual predictor constructs to the explanatory power of an endogenous variable. Effect size f^2 quantifies the change in R^2 when a specific exogenous construct is included in or excluded from the structural model (Cohen, 1988).

$$f^2 = \frac{R^2(\text{included}) - R^2(\text{excluded})}{1 - R^2(\text{included})}$$

Equation 2: Effect size (f^2)

According to established guidelines, f^2 values of 0.02, 0.15, and 0.35 indicate small, medium, and large effects, respectively (Cohen, 1988). In the context of PLS-SEM, effect size is particularly informative in mediation and moderation analyses, as it enables researchers to evaluate the substantive importance of predictors beyond mere statistical significance (Memon et al., 2018; Hair et al., 2021).

By jointly considering path significance, R^2 , and f^2 values, researchers can achieve a nuanced, predictive-oriented, and theoretically grounded interpretation of the structural model results.

3.2.6. Higher-Order Construct

An important feature of the PLS-SEM framework is the possibility to model multidimensional or higher-order constructs (HOCs). Higher-order constructs are abstract theoretical concepts whose indicators are themselves latent constructs, referred to as lower-order constructs (LOCs). Through this hierarchical structure, complex phenomena can be represented at a higher level of abstraction while preserving their multidimensional nature (Sarstedt et al., 2019).

The use of higher-order constructs is particularly valuable when the theoretical concept under investigation encompasses multiple interrelated dimensions that jointly define a broader domain. Instead of modelling each dimension separately in the structural model, higher-order constructs allow researchers to aggregate related dimensions into a single construct, thereby improving model parsimony and interpretability.

3.2.6.1. Advantages of higher-order constructs

The adoption of higher-order constructs offers several methodological and substantive advantages. First, higher-order constructs contribute to model simplification by reducing the number of relationships in the structural model. By replacing multiple lower-order constructs with a single higher-order construct, the complexity of the model is reduced, facilitating interpretation and communication of results (Sarstedt et al., 2019).

Second, higher-order constructs typically exhibit a broader conceptual bandwidth than their lower-order dimensions. Because they capture a wider range of related facets, higher-order constructs often function as stronger predictors in the structural model, leading to increased explanatory power for endogenous constructs (Hair et al., 2019).

Third, the use of higher-order constructs can help mitigate collinearity issues among lower-order dimensions, particularly when these dimensions are highly correlated. By aggregating related constructs at a higher level, redundancy among predictors is reduced, contributing to more stable and reliable parameter estimates.

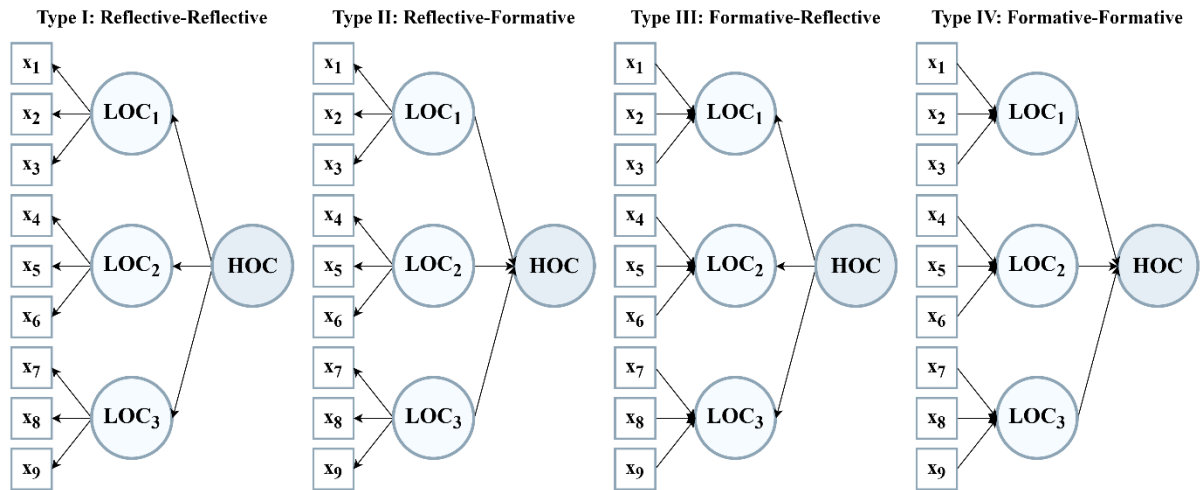
3.2.6.2. Types of higher-order construct models

Higher-order constructs can be specified through different combinations of measurement relationships between constructs and indicators at the lower and higher levels. Importantly, the nature of the measurement relationships at the lower-order level does not necessarily need to coincide with those at the higher-order level. Four main types of higher-order construct models can be distinguished, depending on whether the relationships are reflective or formative at each level:

- **Reflective-Reflective (Type I):** both the lower-order constructs and the higher-order construct are modelled reflectively

- **Reflective-Formative (Type II):** lower-order constructs are reflective, while the higher-order construct is formative
- **Formative-Reflective (Type III):** lower-order constructs are formative, while the higher-order construct is reflective
- **Formative-Formative (Type IV):** both levels are specified formatively

Figure 3.7: Different types of higher-order constructs (adapted from J. F. Hair et al., 2021)



Note: LOC = lower-order component; HOC = higher-order component

The selection of the appropriate higher-order construct type must be grounded in strong theoretical reasoning, as misspecification can lead to biased results and incorrect substantive interpretations.

To implement higher-order constructs in PLS-SEM, researchers can choose among different modelling approaches, the most common being the repeated indicators approach and the two-stage approach (Wetzels et al., 2009).

In the repeated indicators approach, all indicators of the lower-order constructs are assigned simultaneously to the higher-order construct. While this approach is straightforward to implement, it may lead to estimation issues in complex models, particularly when the number of indicators is large or when formative relationships are involved.

Within the two-stage framework, two variants are distinguished: the embedded (or integrated) and the disjoint approach (Sarstedt et al., 2019). The embedded two-stage approach first applies the repeated indicators method to estimate the latent variable scores of the LOCs, which are then used as indicators for the HOC in the second stage. In contrast, the disjoint two-stage approach, the method adopted in the present study, modifies the first stage by omitting the HOC entirely. Instead, the LOCs are directly connected to all other constructs in the model as proxies for the HOC. This allows the latent variable scores of the LOCs to be estimated using their standard

multi-item measures within the full structural context. In the second stage, these scores are then used to estimate the higher-order model. The disjoint approach is preferred here as it provides more stable and unbiased estimates, particularly when the HOC is formative or embedded in a complex network of relationships (Sarstedt et al., 2019). Both variants often yield comparable results, the disjoint approach ensures a more rigorous separation between the measurement of lower-order dimensions and the estimation of the higher-order structure.

3.2.6.3. Methodological considerations

Despite their advantages, higher-order constructs require careful implementation. A frequent methodological error observed in the literature concerns the inadequate assessment of lower-order constructs prior to estimating the higher-order construct. Specifically, reliability and validity of the lower-order measurement models must be established before proceeding to the estimation of the higher-order construct (Sarstedt et al., 2019).

Another common issue arises when researchers incorrectly interpret the relationships between lower-order constructs and the higher-order construct as structural paths rather than as measurement relationships. Such misinterpretations compromise the conceptual clarity of the model and may lead to invalid conclusions. For these reasons, the use of higher-order constructs requires strong theoretical grounding, careful model specification, and rigorous evaluation at both the lower and higher levels of abstraction.

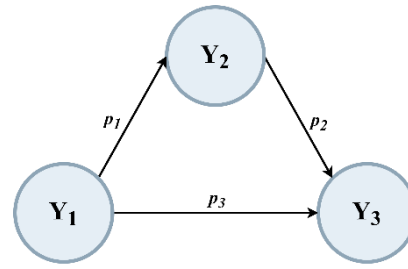
3.2.7. Mediation analysis

Mediation analysis is a core procedure within Structural Equation Modelling and is central to investigating indirect effects among constructs. Rather than examining only direct relationships, mediation analysis enables researchers to test whether the effect of an exogenous construct on an endogenous construct is transmitted through one or more intervening variables, referred to as mediators, thereby providing insight into the underlying process mechanism (Zhao et al., 2010). In PLS-SEM, mediation analysis is especially relevant because complex causal chains can be estimated within a single model while maintaining consistency with the method's variance-based, predictive-oriented rationale (Hair et al., 2021).

3.2.7.1. Conceptualization of mediation effects

A mediation relationship involves three components: an exogenous construct (Y_1), an endogenous construct (Y_3), and a mediator construct (Y_2). The mediator is theorized to account, at least partially, for how and why Y_1 affects Y_3 . In a simple mediation framework, the relationship between Y_1 and Y_3 can be decomposed into a direct effect ($Y_1 \rightarrow Y_3$) and an indirect effect operating through the mediator ($Y_1 \rightarrow Y_2 \rightarrow Y_3$) (Hair et al., 2021).

Figure 3.8: Conceptual representation of a mediation model (adapted from J. F. Hair et al., 2021)



Within this framework, the indirect effect is computed as the product of the path from Y_1 to Y_2 and the path from Y_2 to Y_3 , whereas the total effect corresponds to the sum of the direct and indirect effects and represents the overall influence of Y_1 on Y_3 (Hair et al., 2021).

$$p_1 \cdot p_2$$

Equation 3: Mediation analysis - Indirect effect

$$p_1 \cdot p_2 + p_3$$

Equation 4: Mediation analysis - Total effect

In PLS-SEM, mediation effects are estimated directly as part of the structural model, enabling the simultaneous evaluation of all relevant relationships. Given that the sampling distribution of the product term underlying the indirect effect is often non-normal, the assessment of mediation in PLS-SEM is typically based on bootstrapped confidence intervals, and, where available, bias-corrected confidence intervals are preferred because they provide more accurate coverage in finite samples (Hair et al., 2021). An indirect effect is considered statistically significant when the corresponding (bias-corrected) bootstrapped CI does not include zero.

From a reporting perspective, methodological guidance recommends presenting the magnitude of the indirect effect alongside its confidence interval, and interpreting mediation within the broader model context rather than relying solely on dichotomous significance testing (Memon et al., 2018).

In a complex structural framework such as the one proposed in this study, mediation effects can involve more than one intervening construct and therefore take the form of serial mediation. For instance, Smartphone Addiction may influence Cyber Victimization Risk through a sequential chain such as $SA \rightarrow \text{Gambling} \rightarrow \text{L-RAT} \rightarrow \text{Cyber Victimization Risk}$, where Gambling and L-RAT operate as consecutive mediators transmitting the effect toward the outcome. The distinction between single, multiple, and serial mediation is particularly emphasized in regression-based approaches, where mediation paths are often estimated through sets of separate regression equations requiring researchers to explicitly specify the mediation structure due to the non-simultaneous nature of parameter estimation (Hayes, 2017; Preacher & Hayes, 2008). In contrast, within the PLS-SEM framework, mediation is assessed directly within the structural model and the focus is placed on estimating and testing specific indirect effects (including serial paths) through bootstrapping, without requiring a

formal distinction between single versus multiple mediation or the estimation of separate models for each causal chain (Hair et al., 2021; Memon et al., 2018).

3.2.7.2. Types of mediation

The interpretation of mediation effects is based on the joint evaluation of the direct effect, the indirect effect, and their statistical significance. Depending on the presence, direction, and magnitude of these effects, different mediation scenarios can be identified.

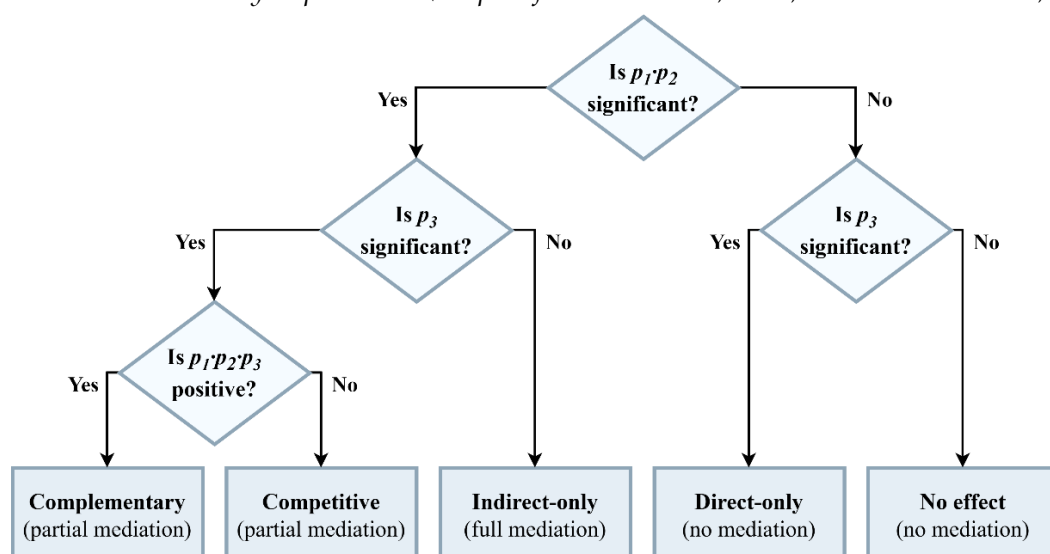
First, **no mediation** is observed when the indirect effect is not statistically significant, regardless of the significance of the direct effect. In this case, the mediator does not transmit the effect of the exogenous construct to the endogenous construct.

Second, **indirect-only mediation (full mediation)** occurs when the indirect effect is statistically significant while the direct effect is not. This pattern indicates that the influence of the exogenous construct on the endogenous construct is entirely transmitted through the mediator.

Third, **partial mediation** is observed when both the indirect and direct effects are statistically significant. Within this category, two distinct situations can be distinguished based on the direction of the effects. Complementary mediation occurs when the direct and indirect effects have the same sign, indicating that both effects reinforce each other. In contrast, competitive mediation arises when the direct and indirect effects have opposite signs, suggesting the presence of countervailing mechanisms influencing the outcome.

Finally, **direct-only (non-mediation)** is identified when the direct effect is statistically significant but the indirect effect is not, indicating that the mediator does not play a meaningful role in the relationship between the exogenous and endogenous constructs.

Figure 3.9: Mediation analysis procedure (adapted from Hair et al., 2021, based on Zhao et al., 2010)



This classification follows the mediation typology and decision logic proposed by Zhao et al. (2010), which emphasizes the central role of the indirect effect and moves beyond the traditional full-versus-partial mediation dichotomy.

3.2.7.3. Methodological considerations in mediation analysis

Several methodological considerations should guide mediation analysis in PLS-SEM. First, mediation hypotheses must be theory-driven: statistical evidence of an indirect effect does not, by itself, establish a causal mechanism. The conceptual ordering of Y_1 , Y_2 and Y_3 must be justified by theory and the research design, including temporal logic where relevant (Hair et al., 2021).

Second, the interpretation of mediation should consider the substantive relevance of the indirect effect. Even statistically significant indirect effects may be practically negligible, particularly in large samples. For this reason, methodological recommendations stress evaluating the magnitude of indirect effects and, where appropriate, relating them to the total effect to provide a more meaningful interpretation of the mediation process (Memon et al., 2018; Hair et al., 2021).

Finally, mediation analysis should be interpreted within the overall structural model, considering explanatory power and theoretical coherence. When properly specified and evaluated, mediation analysis in PLS-SEM provides a rigorous approach for uncovering and quantifying the mechanisms underlying complex relationships among latent constructs.

3.2.8. Moderation analysis

Moderation analysis is employed to examine whether the strength or direction of the relationship between an independent construct and a dependent construct varies as a function of a third variable, referred to as a moderator. While mediation analysis focuses on explaining how or why an effect occurs through intervening mechanisms, moderation analysis addresses when or under what conditions a given relationship changes.

Within the PLS-SEM framework, moderation analysis represents an extension of the structural model and is used to test interaction effects between latent constructs. These effects are specified a priori by introducing an interaction term between the independent construct and the moderator, and they require strong theoretical justification, as they imply conditional hypotheses regarding the variability of structural relationships across levels of the moderator (Hair et al., 2021).

However, the theoretical model proposed in the present study does not include hypotheses involving conditional or interaction effects among constructs. All relationships are specified either as direct effects or as mediated effects, reflecting a process-oriented perspective rather than a conditional one. Moreover, demographic variables such as gender, age, and education are included as control variables with

direct effects, and are not theorized to moderate the structural relationships under investigation.

Accordingly, and in line with methodological recommendations for PLS-SEM, moderation analysis is not performed in this study, as its application would not be theoretically grounded within the proposed framework (Hair et al., 2021; Memon et al., 2018).

3.2.9. Summary of PLS-SEM evaluation procedures and criteria

To enhance transparency and facilitate the interpretation of the empirical results, Table(s) 3.1 - 3.5 summarize the main evaluation procedures and criteria adopted in the PLS-SEM analysis. The tables provide an overview of the constructs assessed, the statistical methods applied, and the threshold values used to evaluate measurement and structural model quality. This summary serves as a methodological guide for the subsequent chapters dedicated to data description and results.

Table 3.1: *Evaluation of reflective measurement models*

Aspect evaluated	What is tested	Method / Indicator	Thresholds / Criteria
Indicator reliability	Strength of indicator-construct relationship	Outer loadings	≥ 0.708 (acceptable ≥ 0.40 if theoretically justified)
Internal consistency reliability	Consistency of indicators measuring the construct	Cronbach's Alpha	≥ 0.70 (conservative lower bound)
		rhoA	≥ 0.70 (most accurate reliability measure)
		Composite Reliability (rhoC)	0.70 - 0.95 (values > 0.95 indicate redundancy)
Convergent validity	Amount of variance captured by the construct	Average Variance Extracted (AVE)	≥ 0.50
Discriminant validity	Distinctiveness of constructs	HTMT ratio	< 0.85 (strict) < 0.90 (lenient)
	Shared variance between constructs	Fornell-Larcker criterion	$\sqrt{\text{AVE}} > \text{inter-construct correlations}$

Table 3.2: *Evaluation of formative measurement models*

Aspect evaluated	What is tested	Method / Indicator	Thresholds / Criteria
Convergent validity	Degree to which formative construct captures the concept	Redundancy analysis	Path coefficient ≥ 0.708
Indicator collinearity	Redundancy among formative indicators	VIF	< 5 (preferably < 3)
Indicator relevance	Relative contribution of indicators	Indicator weights (bootstrap CI)	CI does not include 0
Indicator importance	Absolute contribution of indicators	Indicator loadings	Examined when indicator weights are non-significant

Table 3.3: *Evaluation of the structural model*

Aspect evaluated	What is tested	Method / Indicator	Thresholds / Criteria
Collinearity	Redundancy among predictor constructs	VIF (latent variables)	< 5 (preferably < 3)
Path significance	Strength and direction of relationships	Path coefficients (bootstrapping)	$t > 1.96$ ($\alpha = .05$)
Explanatory power	Variance explained in endogenous constructs	Coefficient of determination (R^2)	0.75 (substantial) 0.50 (moderate), 0.25 (weak)
	Contribution of each predictor	Effect size (f^2)	0.02 (small) 0.15 (medium) 0.35 (large)

Table 3.4: *Evaluation of higher-order constructs (HOC)*

Aspect evaluated	What is tested	Method / Indicator	Thresholds / Criteria
Measurement quality	Validity of lower-order constructs	LOC assessment	Same criteria as first-order constructs
Collinearity	Redundancy among dimensions	VIF	< 5
Path significance	Contribution of dimensions to HOC	Bootstrapped paths	CI does not include 0

Table 3.5: *Mediation analysis*

Aspect evaluated	What is tested	Method / Indicator	Thresholds / Criteria
Indirect effects	Presence of mediation	Bootstrapped indirect effects	CI (bias-corrected) excludes 0
Type of mediation	Nature of mediation	Direct vs indirect effects	Zhao et al. (2010) classification
Substantive relevance	Importance of mediation	Magnitude of indirect effect	Interpreted relative to total effect

3.3. Machine Learning

3.3.1. Machine Learning and predictive modelling

While Partial Least Squares Structural Equation Modelling (PLS-SEM) represents the primary analytical framework adopted in this study to test theoretically grounded relationships among latent constructs, it is complemented by Machine Learning (ML) techniques to address a distinct but related research objective, namely predictive accuracy. The joint use of PLS-SEM and ML reflects the increasing recognition in methodological literature that explanatory and predictive modelling serve different

purposes and should be viewed as complementary rather than mutually exclusive approaches (Hair et al., 2021; Shmueli et al., 2019).

PLS-SEM is particularly suited for theory testing and development, as it allows researchers to model complex causal structures, account for measurement error, and assess mediation mechanisms within a single coherent framework. However, its primary focus lies on explaining variance and evaluating hypothesized relationships, rather than optimizing predictive performance on unseen data. As a consequence, models that exhibit strong explanatory power do not necessarily achieve high levels of out-of-sample prediction accuracy. Machine Learning approaches, by contrast, are explicitly designed to maximize predictive performance. ML models prioritize generalization to new data, often relying on algorithmic learning strategies that can capture non-linear relationships and complex interactions among predictors without imposing strong parametric assumptions. In the context of this study, ML is therefore employed to complement the explanatory insights derived from PLS-SEM by evaluating the extent to which the observed patterns can be leveraged to accurately classify individuals with respect to cyber victimization risk, defined as a binary outcome.

The rationale for integrating ML into the analytical strategy is thus not to replace PLS-SEM, but to extend the analysis toward a predictive perspective, allowing for a more comprehensive assessment of the phenomenon under investigation. This dual approach aligns with recent methodological recommendations advocating the joint use of explanatory and predictive techniques in behavioural and social science research (Hair et al., 2021). To clarify the distinct roles and underlying assumptions of PLS-SEM and Machine Learning, Table 3.6 summarizes their main differences along key analytical dimensions.

Table 3.6: Comparison between PLS-SEM and Machine Learning approaches

Dimension	PLS-SEM	Machine Learning
Primary objective	Explanatory / confirmatory	Predictive / verificative
Relationship specification	Explicitly specified by theory	Learned from data
Assumptions	Relatively strong (e.g.: model structure)	Minimal distributional assumptions
Ability to test mediation	High	Limited / indirect
Handling of non-linearity	Limited	High
Interpretability	High	Moderate
Evaluation focus	Path significance, R^2 , effect sizes	Predictive performance (e.g.: accuracy, AUC)
Use in this study	Testing theoretical relationships	Robustness and variables ranking

By combining PLS-SEM and Machine Learning, the present study adopts a methodological strategy that leverages the strengths of both approaches: PLS-SEM provides a theoretically grounded explanation of the relationships among smartphone addiction, gambling behaviour, routine activities, and cyber victimization risk, while Machine Learning, through supervised classification models, assesses the practical predictive value of these relationships when applied to individual-level data.

This integrated framework allows the study to address not only why and how cyber victimization risk emerges, but also how accurately such risk can be predicted, thereby enhancing both the theoretical and applied relevance of the findings.

3.3.2. Machine Learning paradigms and supervised learning models

Machine Learning (ML) can be defined as a set of computational methods that use algorithms to identify patterns in data and improve their performance on a specific task through experience (Mitchell, 1997). Unlike traditional statistical approaches that rely on pre-defined functional forms, ML systems are designed to "learn" directly from the data, optimizing a performance measure by minimizing prediction error or maximizing classification accuracy (James et al., 2021). Depending on the nature of the learning process and the availability of labelled data, ML paradigms are categorized into three main types: supervised, unsupervised, and reinforcement learning (Hastie et al., 2009).

Supervised learning involves training an algorithm on a labelled dataset, where each observation consists of an input vector of predictors (features) and a corresponding output (target). The model learns a general rule or mapping function to predict the target for new, unseen data. These models are primarily used for classification (categorical outcomes) and regression (continuous outcomes) tasks. Unsupervised learning is applied when the data lacks labels. The objective is to discover hidden structures, patterns, or clusters within the data without a pre-defined outcome variable. Common applications include customer segmentation through clustering or feature reduction through Principal Component Analysis (PCA). Reinforcement learning involves an autonomous agent that learns to make sequences of decisions by interacting with an environment. The agent receives feedback in the form of rewards or penalties and aims to maximize the cumulative reward over time. While powerful for robotics and gaming, this paradigm is less common in social science research focused on static survey data.

Given that the objective of the present study is to predict a specific and labelled outcome (i.e., cyber victimization risk), supervised learning is the most appropriate paradigm. Within this framework, several algorithms exist, each with distinct mechanisms:

- **Logistic Regression:** A baseline model for binary classification that estimates the probability of an observation belonging to a specific class

using the logistic function. It is highly interpretable but limited by the assumption of linearity between predictors and the log-odds of the outcome.

- **k-Nearest Neighbours (k-NN):** A non-parametric, instance-based method that classifies an observation based on the majority class of its closest neighbours in the feature space. While intuitive, its performance can degrade with high-dimensional data or irrelevant features.
- **Decision Trees:** These models use a tree-like structure to recursively partition the data into subsets (nodes) based on the value of predictors, aiming for maximum homogeneity in the resulting groups. They are easy to visualize but highly prone to overfitting, capturing noise instead of the underlying signal.
- **Support Vector Machines (SVM):** This algorithm identifies the optimal hyperplane that separates classes with the maximum margin. By using "kernels" SVM can transform data into higher dimensions to handle non-linear separations, though it can be computationally intensive and less transparent than simpler models.
- **Ensemble Methods:** These techniques combine multiple individual learners to produce a more robust and accurate model. Bagging (Bootstrap Aggregating) reduces variance by averaging predictions from multiple independent models, while Boosting sequentially builds models that correct the errors of their predecessors.

Among these, the Random Forest algorithm represents a superior choice for this study. Random Forest is an ensemble method based on bagging that grows a large number of independent decision trees. Each tree is trained on a random bootstrap sample of the data, and at each split, only a random subset of predictors is considered (Breiman, 2001). The final classification is determined through majority voting across all trees.

From a methodological perspective, several characteristics motivate the adoption of Random Forest in this context. First, Random Forest is capable of capturing non-linear relationships and complex interaction effects among predictors without requiring explicit model specification. This property is especially relevant in behavioural and social science applications, where relationships between individual characteristics and outcomes are rarely purely linear.

Second, Random Forest is known for its robustness to overfitting, particularly when compared to single decision trees. By combining multiple weak learners and introducing randomness both in data sampling and feature selection, the algorithm reduces variance and improves generalization performance on unseen data.

Third, Random Forest performs well in settings characterized by correlated predictors and mixed types of input variables, which are common in survey-based research involving psychological, behavioural, and socio-demographic measures. Unlike parametric models, it does not rely on strict distributional assumptions and is relatively insensitive to multicollinearity.

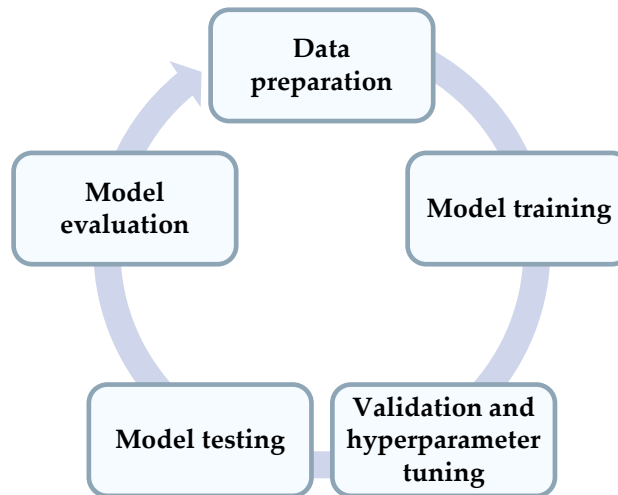
Finally, Random Forest provides measures of variable importance, which allow researchers to assess the relative contribution of individual predictors to the classification task. Although Machine Learning models are primarily predictive in nature, this feature facilitates a meaningful dialogue with the results obtained from PLS-SEM by highlighting which variables play a central role in predicting cyber victimization risk.

By integrating the explanatory depth of PLS-SEM with the predictive robustness of Random Forest, this study ensures that the findings regarding cyber victimization are both theoretically grounded and practically applicable for risk classification (James et al., 2021).

3.3.3. Machine Learning lifecycle

Machine Learning models are developed and evaluated through a structured and iterative process commonly referred to as the Machine Learning lifecycle. This lifecycle provides a systematic framework that guides the progression from raw data to validated predictive models and is essential for ensuring robust performance, reproducibility, and proper generalization to unseen data. Adhering to a clearly defined lifecycle is particularly important in supervised learning contexts, where predictive accuracy and avoidance of data leakage are central methodological concerns (Boehmke & Greenwell, 2019; Hastie et al., 2009). The main phases of the lifecycle include data preparation, model training, validation and hyperparameter tuning, model testing, and performance evaluation.

The first phase of the Machine Learning lifecycle involves data preparation, which constitutes a critical step in determining the quality and reliability of the predictive model. This phase includes the selection of predictor variables, handling of missing data, and appropriate transformation or encoding of features when required. In supervised learning tasks, particular attention must be paid to the correct definition of the outcome variable and to ensuring that all predictors reflect information that would be available at the time prediction is made. Unlike structural equation modelling approaches, Machine Learning algorithms operate directly on observed variables and do not explicitly model measurement error. As a consequence, careful preprocessing is necessary to reduce noise, inconsistencies, and potential sources of bias in the data. In this study, data preparation is performed prior to model estimation to ensure that the Random Forest algorithm is trained on a consistent and well-defined feature set.

Figure 3.10: Main phases of the Machine Learning lifecycle

Following data preparation, the Machine Learning model is trained on a subset of the available observations. During training, the algorithm learns the relationship between the predictor variables and the binary outcome by minimizing classification error on the training data. In the case of Random Forest, this process involves fitting a large number of decision trees to bootstrap samples of the training set, with randomness introduced both in the selection of observations and in the subset of features considered at each split. The objective of the training phase is not to achieve perfect fit to the training data, but to learn patterns that generalize beyond the observed sample. Models that adapt too closely to the training data may exhibit overfitting and consequently perform poorly when applied to new observations.

The validation phase is used to optimize model performance through hyperparameter tuning. Hyperparameters are configuration parameters that control the learning process but are not estimated directly from the data, such as the number of trees or the number of predictors considered at each split in a Random Forest model. Validation is typically implemented using resampling techniques such as cross-validation, which allow the model to be evaluated repeatedly on different subsets of the data. This process supports the identification of hyperparameter configurations that balance predictive accuracy and model stability, thereby reducing the risk of overfitting (Hastie et al., 2009).

Once the optimal hyperparameter configuration has been selected, the model is evaluated on a test set that has not been used during either training or validation. The test phase provides an unbiased estimate of the model's predictive performance and represents the most realistic assessment of how the model would perform on new, unseen data. A strict separation between training, validation, and testing data is a fundamental principle of Machine Learning. Violations of this principle may lead to overly optimistic performance estimates and undermine the credibility of the predictive analysis (Boehmke & Greenwell, 2019).

The final phase of the Machine Learning lifecycle concerns model evaluation, which focuses on assessing predictive performance using appropriate metrics. In supervised classification problems with a binary outcome, such as the prediction of cyber victimization risk, evaluation typically relies on metrics that capture both overall classification accuracy and the model's ability to discriminate between classes.

In addition to predictive performance, Machine Learning models such as Random Forest allow for the computation of variable importance measures, which provide insight into the relative contribution of individual predictors to the classification task. While these measures do not imply causal relationships, they offer a valuable descriptive complement to the explanatory findings derived from PLS-SEM.

Overall, the Machine Learning lifecycle provides a rigorous and transparent framework for predictive modelling. By explicitly following each phase of the lifecycle, the present study ensures that the Random Forest analysis adheres to established best practices in Machine Learning and yields reliable and interpretable predictive results.

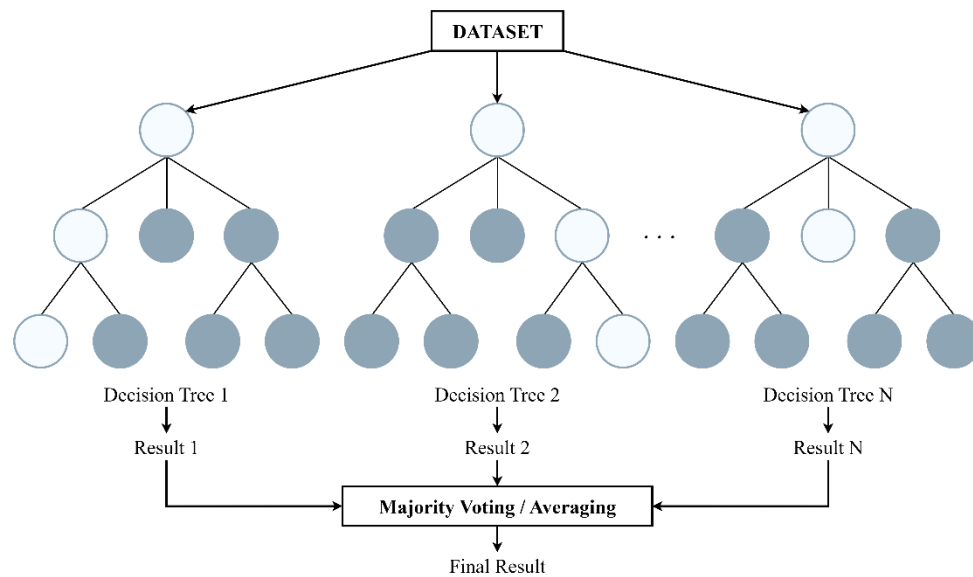
3.3.4. Random Forest

Random Forest is a supervised Machine Learning algorithm belonging to the family of ensemble methods, specifically designed to improve predictive performance and robustness by combining the outputs of multiple decision trees. Originally introduced by Breiman (2001), Random Forest has become one of the most widely used algorithms for classification and regression tasks, particularly in applied research contexts characterized by complex, non-linear relationships and correlated predictors.

3.3.4.1. The algorithm

At its core, Random Forest operates by constructing a large collection of decision trees and aggregating their predictions to produce a final output. In classification tasks, such as the binary prediction of cyber victimization risk, the final class assignment is typically determined by majority voting across all trees in the forest.

Each individual decision tree partitions the feature space through a sequence of binary splits, selected to maximize node purity according to criteria such as the Gini index or entropy. While single decision trees are highly flexible and capable of modelling complex interactions, they are also prone to instability and overfitting. Random Forest addresses these limitations by introducing randomness and aggregation mechanisms that stabilize predictions and enhance generalization performance (Breiman, 2001; Hastie et al., 2009).

Figure 3.11: Conceptual representation of the Random Forest algorithm

3.3.4.2. Ensemble learning, bagging, and feature randomness

Random Forest belongs to the class of ensemble learning methods, which combine multiple base learners to obtain a predictive model that outperforms individual components. The ensemble strategy employed by Random Forest is based on bootstrap aggregation (bagging), whereby each decision tree is trained on a bootstrap sample drawn with replacement from the original training dataset.

Bagging reduces model variance by averaging over multiple models trained on slightly different datasets. However, if all trees are highly correlated, the variance reduction achieved through aggregation may be limited. To further decorrelate the trees, Random Forest introduces an additional source of randomness through feature subsampling: at each split, only a random subset of predictors is considered as candidates for splitting.

This combination of bagging and feature randomness ensures that individual trees explore different regions of the feature space and rely on different subsets of predictors, thereby increasing diversity within the ensemble. The resulting reduction in correlation among trees is a key factor underlying the strong predictive performance of Random Forest (Breiman, 2001; Hastie et al., 2009).

Furthermore, the bagging procedure allows for the estimation of the model's generalization error through the Out-of-Bag (OOB) error. Since each tree is trained on approximately two-thirds of the original data (the bootstrap sample), the remaining one-third of observations, which are not used for that specific tree, act as an internal and unbiased test set. This mechanism provides a continuous and robust performance evaluation during the training process, eliminating the need to hold out additional data from the training set for validation purposes (Breiman, 2001).

3.3.4.3. Comparison with single decision trees

Compared to single decision trees, Random Forest offers several methodological advantages. First, while individual trees tend to exhibit low bias but high variance, Random Forest substantially reduces variance through aggregation without significantly increasing bias. This leads to more stable and reliable predictions across different samples.

Second, Random Forest is less sensitive to small perturbations in the data. Whereas minor changes in the training set can lead to substantially different tree structures in single decision trees, the ensemble nature of Random Forest smooths out such instabilities. This property is particularly valuable in social and behavioural datasets, which often contain noise and measurement variability.

Finally, Random Forest generally outperforms single decision trees in terms of predictive accuracy, especially in settings involving non-linear effects, high-order interactions, and correlated predictors, conditions that are common in empirical studies based on survey data (Hastie et al., 2009; Boehmke & Greenwell, 2019).

3.3.4.4. Hyperparameters of Random Forest

The performance of a Random Forest model depends on several hyperparameters that control model complexity and learning behaviour. Among the most important hyperparameters are:

- the **number of trees** in the forest, which determines the size of the ensemble;
- the number of predictors considered at each split (commonly referred to as *mtry*);
- the **maximum depth** of the trees or, alternatively, the minimum number of observations required to split a node;
- the **minimum node size**, which controls how finely the feature space is partitioned.

Increasing the number of trees generally improves predictive performance by stabilizing the ensemble, although gains tend to plateau beyond a certain point. The choice of *mtry* plays a critical role in balancing tree diversity and individual tree strength: smaller values increase randomness and reduce correlation among trees, while larger values allow trees to exploit stronger predictors more consistently (Breiman, 2001).

Unlike other ensemble methods (such as Gradient Boosting), Random Forest possesses the distinctive property of not overfitting as the number of trees increases. Although an excessive number of trees increases computational cost, the generalization error does not worsen; instead, it tends to stabilize or plateau. This

makes the algorithm extremely robust and less sensitive to the precise selection of this specific hyperparameter (Hastie et al., 2009).

Hyperparameter tuning is therefore essential to achieving optimal model performance and is typically conducted using resampling-based validation techniques, as discussed in the Machine Learning lifecycle.

3.3.4.5. Bias - variance trade-off in Random Forest

The effectiveness of Random Forest can be understood through the lens of the bias - variance trade-off, a central concept in Machine Learning and Statistical Learning theory. Decision trees are low-bias, high-variance models: they can closely fit training data but often generalize poorly. Random Forest mitigates this issue by averaging over many trees, thereby substantially reducing variance.

While the introduction of randomness may slightly increase bias compared to a fully optimized single tree, the overall reduction in variance typically leads to a net improvement in generalization performance. As a result, Random Forest often achieves a favourable balance between bias and variance, making it particularly well suited for predictive tasks in complex, real-world datasets (Hastie et al., 2009; Boehmke & Greenwell, 2019).

In summary, Random Forest combines flexible modelling capacity with strong regularization through ensemble learning, offering a robust and theoretically grounded approach to supervised classification. These properties justify its adoption in the present study as a complementary predictive tool alongside explanatory modelling techniques.

3.3.5. Evaluation of the Machine Learning model

The evaluation of a Machine Learning model represents a fundamental step in assessing its ability to generalize to new, unseen data and to provide reliable predictions. In supervised classification tasks, model evaluation focuses on quantifying predictive performance using metrics that capture both overall classification accuracy and the model's capacity to discriminate between outcome classes. In the present study, the Random Forest model is evaluated with respect to its ability to predict cyber victimization risk, operationalized as a binary outcome.

Consistent with best practices in Machine Learning, model evaluation is conducted on a test dataset that is not used during model training or hyperparameter tuning. This strict separation between training and testing phases ensures unbiased performance estimates and guards against optimistic evaluation due to overfitting (Boehmke & Greenwell, 2019).

3.3.5.1. Classification performance metrics and interpretation

Several complementary metrics are employed to evaluate the performance of the Random Forest classifier. Each metric captures a different aspect of predictive performance and contributes to a comprehensive assessment of model quality.

First, a confusion matrix is used to provide a detailed representation of classification outcomes, enabling a detailed interpretation of model errors and facilitating the assessment of the practical implications of false positive and false negative predictions.

Figure 3.12: Confusion matrix and model's performance evaluation metrics (adapted from Encord, n.d.)

		Predicted class		
		Positive	Negative	
Actual class	Positive	True Positive (TP)	False Negative (FN) Type II Error	Sensitivity $\frac{TP}{TP + FN}$
	Negative	False Positive (FP) Type I Error	True Negative (TN)	Specificity $\frac{TN}{TN + FP}$
		Precision $\frac{TP}{TP + FP}$	Negative Predictive Value $\frac{TN}{TN + FN}$	Accuracy $\frac{TP + TN}{TP + TN + FP + FN}$

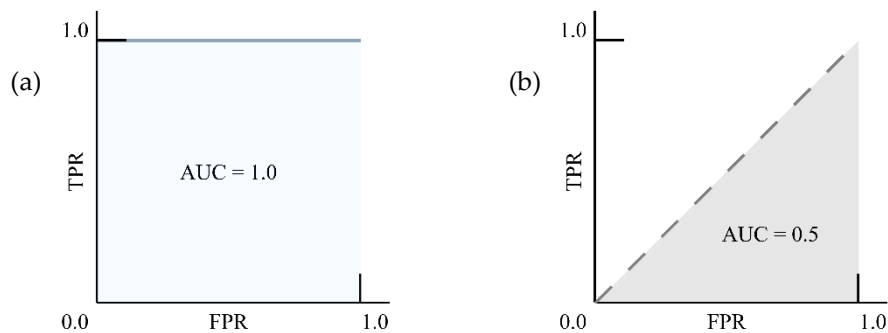
A first important metric is Accuracy, that measures the proportion of correctly classified observations relative to the total number of cases. Although accuracy provides an intuitive measure of overall performance, it may be misleading when class distributions are unbalanced (Encord, n.d.).

To complement accuracy, the Receiver Operating Characteristic (ROC) curve and the corresponding Area Under the Curve (AUC) are examined. The ROC curve plots the True Positive Rate (sensitivity) against the False Positive Rate across different classification thresholds, while the AUC measures the model's ability to assign higher predicted probabilities to the positive class than to the negative class.

$$(a) TPR = \frac{TP}{TP + FN} \quad (b) FPR = \frac{FP}{FP + TN}$$

Equation 5: True Positive Rate (a) and False Positive Rate (b)

AUC values range from 0.5 (no discriminative ability) to 1 (perfect discrimination), with higher values indicating stronger classification performance (Hastie et al., 2009).

Figure 3.13: ROC and AUC of a perfect model (a) and random guessing (b)

While accuracy and AUC provide a broad overview of model performance, classification quality is further scrutinized through Precision, Recall (Sensitivity), and the F1-score. Precision measures the model's ability to correctly identify only truly positive cases, minimizing false alarms. Recall assesses the model's capacity to capture all actual positive instances, which is critical in risk-identification contexts. The F1-score, representing the harmonic mean of Precision and Recall, offers a balanced performance measure that is particularly robust when dealing with potential class imbalances (Hastie et al., 2009). In many classification settings, Precision and Recall exhibit an inherent trade-off: improving the model's ability to capture all true positive cases (high Recall) often comes at the cost of generating more false alarms (lower Precision), whereas prioritizing Precision may reduce the model's sensitivity to at-risk individuals. Evaluating both metrics jointly is therefore essential for balancing missed detections and false positives in risk prediction contexts.

In binary classification problems, predicted probabilities must be converted into class labels by selecting a classification threshold. While a threshold of 0.5 is commonly adopted, alternative thresholds may be considered depending on the relative importance of sensitivity and specificity. In the context of cyber victimization risk prediction, threshold selection involves balancing the correct identification of at-risk individuals against the risk of false alarms. Evaluating model performance across multiple thresholds further strengthens the robustness of the predictive assessment and prevents overreliance on arbitrary decision rules.

3.3.5.2. Variable importance analysis

Beyond predictive accuracy, Random Forest models allow for the computation of variable importance measures, which quantify the contribution of individual predictors to the classification task. Common approaches include impurity-based importance and permutation-based importance, both of which assess how strongly each variable influences model predictions.

Although variable importance measures do not imply causal relationships, they provide a valuable descriptive complement to the explanatory findings derived from PLS-SEM. In this study, variable importance analysis is used to identify which

behavioural, psychological, and contextual factors are most influential in predicting cyber victimization risk, thereby supporting an integrated interpretation of predictive and explanatory results (Boehmke & Greenwell, 2019).

The results of this analysis are typically visualized through variable importance plots, which rank predictors based on their relative contribution to model accuracy or node purity (e.g.: Mean Decrease Gini). This graphical representation facilitates a clear comparison between behavioural, psychological, and contextual risk factors, highlighting the primary drivers of cyber victimization risk.

3.3.5.3. Summary of Machine Learning evaluation criteria

To enhance transparency and facilitate interpretation of the predictive results, Table 3.7 summarizes the main evaluation criteria adopted for the Random Forest classification model.

Although Machine Learning evaluation focuses on predictive performance, results must be interpreted with caution. High predictive accuracy does not necessarily imply theoretical validity or causal relevance, just as strong explanatory relationships do not guarantee superior predictive performance. For this reason, the Machine Learning results are interpreted as complementary to, rather than substitutes for, the PLS-SEM findings.

By jointly considering classification performance metrics and variable importance measures, the evaluation of the Random Forest model provides a robust and methodologically sound assessment of predictive capability while maintaining conceptual coherence with the explanatory framework developed in the earlier stages of the analysis.

Table 3.7: Evaluation of the Machine Learning model

Aspect evaluated	What is assessed	Metric / Method	Interpretation
Overall performance	Correct classification rate	Accuracy	Higher values indicate better overall performance
Discriminative ability	Separation between classes	ROC curve, AUC	AUC $\approx$ 0.5 (random), AUC $\rightarrow$ 1 (strong discrimination)
Error structure	Detail of misclassifications	Confusion matrix (Precision, Recall, F1-score)	Balancing false alarms (Precision) vs. missed cases (Recall)
Threshold robustness	Stability across cutoffs	Performance across thresholds	Avoid reliance on a single arbitrary cutoff
Predictor relevance	Contribution of predictors	Variable importance (Importance Plots)	Ranking of key drivers; non-causal interpretation

4 Data, variables, and modelling strategy

This chapter describes the dataset, the operationalization of variables, and the analytical strategies adopted in the empirical analysis. First, data source and sample characteristics are presented, followed by the procedures used for data cleaning and preprocessing. It then details the measurement of the dependent variable, explanatory constructs, and control variables, with explicit reference to the survey items and coding criteria. Finally, the chapter outlines the modelling strategies adopted in the study, distinguishing between the structural equation modelling approach and the machine learning framework.

All analyses and preprocessing steps were implemented in RStudio 2024.12.1 Build 563.

4.1. Dataset description

The analysis is based on data drawn from the National Inquiry on Cybersecurity and Confidence in Spanish Households, conducted by the Spanish Ministry of Economic Affairs and Digital Transformation between January 2019 and June 2021. The survey represents a large-scale, nationally representative data collection effort aimed at analysing cybersecurity incidents, digital behaviours, and psychosocial factors associated with Internet use among Spanish households.

The inquiry follows a semi-annual panel design, with approximately 3,500 participants per wave. Respondents are recruited at the household level, and individuals aged 15 years or older who regularly use the Internet are invited to participate. To preserve representativeness over time, non-respondents are systematically replaced, ensuring stability across key socio-demographic dimensions, including geographical region, urban size, socio-economic status, and household composition.

For the purposes of this study, a subsample of the original database was extracted. Specifically, only respondents who completed the questionnaire at least once between 2019 and 2021 were retained. To avoid duplication and ensure independence of observations, only the first available response for each individual was

selected. This decision allows the analysis to rely on a cross-sectional snapshot per participant, facilitating comparability across constructs and modelling approaches.

Data collected after 2021 were excluded from the analysis due to the loss of unique respondent identifiers, which would prevent reliable de-duplication of observations across waves. Earlier waves were also excluded, as they did not contain several key items required for the operationalisation of the constructs analysed in this study, or presented differences in item wording and scale structure.

The final analytical sample consists of $N = 7,064$ individuals, of whom 777 respondents reported having used online betting or gambling platforms. The remaining participants constitute the non-gambler subgroup. The sample includes respondents who completed the survey via smartphone or tablet ($n = 6,720$), as well as a smaller proportion who completed the questionnaire using a personal computer ($n = 884$).

Overall, the resulting dataset provides a large, heterogeneous, and nationally representative sample, suitable for analysing the relationships between smartphone addiction, social and behavioural factors, lifestyle-routine activity patterns, and cybervictimization risk. The size and structure of the sample meet the methodological requirements for both PLS-SEM estimation and Machine Learning predictive modelling, as discussed in the following sections.

4.2. Data preparation and preprocessing

The analytical dataset used in this thesis derives from the broader national survey database described in Section 4.1. Before model estimation, a structured cleaning and preprocessing workflow was applied to ensure that (i) the observations were coherent with the study's temporal scope, (ii) the indicators selected from the questionnaire were methodologically aligned with the theoretical framework, and (iii) the resulting data matrix satisfied basic requirements for both PLS-SEM estimation and the complementary Machine Learning phase.

4.2.1. Item selection and theoretical alignment

Survey items were selected by taking the January 2021 questionnaire as the operational reference, ensuring consistency in wording, response options, and item identifiers across the measurement model. The selection was further informed by the operationalization strategy adopted in prior research conducted on the same database, in particular the application of Lifestyle-Routine Activity Theory (LRAT) to cybercrime victimization and smartphone-related vulnerabilities (Herrero, Torres, Vivas, Hidalgo, et al., 2021).

This approach is substantively important for two reasons. First, it improves construct comparability with existing empirical evidence, especially for LRAT

components and cybervictimization-related indicators. Second, it reduces the risk of “construct drift” (i.e., measuring conceptually adjacent but not equivalent items), which can undermine both the interpretability of PLS-SEM paths and the validity of predictive comparisons in the Machine Learning phase.

4.2.2. Temporal filtering and removal of invalid observations

Coherently with the inclusion criteria defined in Section 4.1, only valid responses collected between January 2019 and June 2021 were retained. Observations outside this time window were excluded to avoid mixing potentially heterogeneous measurement conditions and to ensure alignment with the target population and study period.

Within the selected time range, the dataset was further cleaned by removing incomplete responses. This step is particularly relevant in survey-based research because partial completion often produces structurally missing blocks of items (not missing-at-random patterns), which cannot be reliably imputed without introducing bias into latent variable scores and derived indices.

4.2.3. Recoding of sociodemographic variables

Several sociodemographic variables were recoded to improve usability in modelling and to avoid unnecessarily sparse categories. Two variables, in particular, required explicit operationalization choices:

- Age: originally spanning from 17 to 75 years, it was grouped into five ordinal age intervals.
- Educational background: the original eight response categories were consolidated into four broader categories.

An ordinal encoding strategy was used where appropriate. The main rationale was methodological and computational: reducing the cardinality of variables that carry limited incremental information improves model parsimony, reduces fragmentation in category-based controls, and supports computational efficiency in predictive pipelines (INRIA, 2025). Importantly, this recoding was applied only to variables where the aggregation does not alter the substantive meaning of the control (i.e., preserving the intended rank/ordering for age and education).

4.2.4. Consistency checks and reverse coding

Before constructing composite scores and estimating latent variables, item directionality was verified across all indicators. Because several scales include negatively worded items, reverse coding was applied when needed so that higher values consistently correspond to higher levels of the underlying construct (e.g.: higher target suitability, higher perceived social digital pressure, higher social support).

This harmonization step is essential in latent variable modelling: failing to reverse-code negatively keyed indicators would induce artificial attenuation in internal consistency metrics and could distort outer loadings/weights, thereby compromising both measurement quality and the interpretation of structural effects.

4.2.5. Data scaling for PLS-SEM and Machine Learning

A key methodological point is that preprocessing was not identical for the explanatory (PLS-SEM) and predictive (Machine Learning) phases, because the two approaches have different requirements and conventions.

For the PLS-SEM phase, no explicit rescaling was required prior to constructing latent variables, as the indicators contributing to each construct share common response formats within the questionnaire and PLS-SEM estimation does not rely on the same standardization assumptions as many parametric regression routines. In practice, the estimation procedure can work directly with the original item scales when they are coherent within constructs and the measurement specification is appropriate (Hair et al., 2021; Haenlein & Kaplan, 2004).

By contrast, for the Machine Learning phase, a dedicated transformation was applied to support comparability and interpretability of the feature space. Specifically, variables used as predictors in ML were transformed using a POMP (Percent of Maximum Possible) scaling approach and shifted to a 0-100 range, enabling more objective comparison across features that may originate from different item groupings or score constructions (Saqr & López-Pernas, 2024). This step is consistent with the broader aim of the ML component in this thesis: evaluating predictive performance on a standardized feature representation rather than relying on the idiosyncrasies of raw questionnaire coding.

4.2.6. Descriptive statistics and reliability checks

Finally, prior to estimation, preliminary diagnostics were run to validate the basic quality of the cleaned dataset:

- Descriptive statistics were computed for all variables retained in the study, supporting an initial assessment of distributions, variability, and potential ceiling/floor effects.
- Reliability-oriented checks were conducted as an initial diagnostic for reflective blocks (where applicable), anticipating the more formal measurement model evaluation procedures presented later in the thesis.

Overall, these cleaning and preprocessing procedures ensured that the final dataset was coherent with the questionnaire structure and the study's theoretical framework, while also satisfying the practical requirements of (i) PLS-SEM estimation

with reflective and formative components and (ii) predictive modelling with standardized feature inputs.

4.3. Variables, constructs, and measurement scales

This section describes the variables included in the empirical analysis and explains how they are operationalized using the survey questionnaire. For each variable and construct, the original survey questions and items are reported, together with information on response formats and coding procedures. This detailed presentation ensures clarity and facilitates replication of the analyses.

4.3.1. Dependent variable: Cybervictimization risk

Cybervictimization risk is the dependent variable of this study. It captures whether respondents experienced cybersecurity incidents during their everyday use of digital devices and online services. The variable is based on self-reported information collected in the Security Incidents module of the questionnaire and refers to events experienced in the six months preceding the survey.

Cybervictimization risk is derived from the following survey question:

“In the last six months, have you experienced any of the following security issues on your Computer/Smartphone/Tablet that you typically use to access the Internet?”

The question includes a list of different types of incidents related to malware, unauthorized access, fraud, and other cybersecurity threats. Of the ten possible victimisation situations, six were selected for this study:

- **pc1a_pcan_1:** Computer viruses or other malicious code/malware (Trojans, worms, etc.)
- **pc1a_pcan_2:** I have lost or deleted data or files
- **pc1a_pcan_4:** Someone has accessed my email or social media profile and impersonated me or created a false identity about me
- **pc1a_pcan_5:** Someone has accessed my computer and/or devices without my consent (intrusion)
- **pc1a_pcan_6:** Loss/theft of my device
- **pc1a_pcan_7:** I have been unable to access an online service due to cyberattacks (downed websites, etc.)

The selection of these six incidents from the broader list available in the questionnaire was guided by psychological and practical considerations, following criteria previously adopted in empirical research using the same dataset (Urueña López et al., 2019). Each item is coded as a dichotomous indicator (0 = no / 1 = yes) reflecting the occurrence of the corresponding incident. These items are then combined

into a single composite indicator, labelled **pc1a_Riesgo**, which represents the overall level of cybervictimization risk experienced by the respondent. This aggregation procedure produces a risk profile ranging from 0 to 6, where 0 indicates that the respondent did not experience any of the selected victimization events in the reference period, and 6 represents the highest level of cybervictimization risk, corresponding to exposure to all six incidents.

The internal consistency of the six dichotomous items was assessed using the Kuder-Richardson Formula 20 (KR-20). Results indicate high reliability, with all items exceeding the recommended threshold of 0.50 for unidimensional analysis. These findings support the use of a single composite indicator, rather than modelling each incident as a separate dichotomous variable, allowing for a more parsimonious and interpretable representation of cybervictimization risk.

In the analytical sample, the average value of the index is 0.37, with a standard deviation of 0.80, indicating a generally low prevalence of victimization events, but with substantial variability across respondents.

Table 4.1: *Cybervictimization risk - Descriptive statistics*

Item	Mean	SD	Median	Min	Max	Skewness	Kurtosis	SE
pc1a_Riesgo	0.37	0.80	0	0	6	2.90	10.61	0.01

4.3.2. L-RAT framework

The L-RAT framework is operationalized through four components reflecting different aspects of users' online routines and security-related behaviours: Exposure to motivated offenders, Proximity to motivated offenders, Target suitability, and Capable guardianship.

4.3.2.1. Exposure to motivated offenders

Exposure to motivated offenders captures the extent to which respondents engage in online activities that may increase contact with potentially harmful or malicious actors. The construct reflects routine participation in digital services and online environments where cybersecurity threats are more likely to occur.

Exposure is measured using items from the Internet services usage module of the questionnaire, referring to services used by the respondent during the previous six months. The construct is derived from the following survey question:

“From the following list of services offered on the Internet, indicate whether you have used them in the last six months from your computer/smartphone/tablet”

Based on this question, the following items are included in the construct:

- **a1_pcan_6:** Online payments (PayPal, Bizum, etc.)
- **a1_pcan_7:** Accessing and/or downloading free digital content (YouTube, Spotify, Vimeo, Telegram, social media, etc.)
- **a1_pcan_8:** Downloading files (programmes, videos, films, music, video games) via P2P networks (eMule, BitTorrent, U Torrent, Elite Torrent, Mejortorrent, tumejortorrent, grantorrent, don torrent, torrentz.eu) or unofficial websites
- **a1_pcan_13:** Visiting adult entertainment sites (e.g.: pornography)

Each item indicates whether the respondent used the corresponding online service during the reference period and is coded as a dichotomous indicator (0 = no / 1 = yes). Mean value of the four items is 0.47 and the standard deviation is 0.45.

Table 4.2: L-RAT framework - Descriptive statistics of Exposure to motivated offender

Item	Mean	SD	Median	Min	Max	Skewness	Kurtosis	SE
Exposure (Average)	0.47	0.45	0	0	1	0.19	-0.90	0.01
a1_pcan_6	0.63	0.48	1	0	1	-0.54	-1.71	0.01
a1_pcan_7	0.74	0.44	1	0	1	-1.10	-0.78	0.01
a1_pcan_8	0.27	0.44	0	0	1	1.06	-0.87	0.01
a1_pcan_13	0.22	0.42	0	0	1	1.33	-0.24	0.00

4.3.2.2. Proximity to motivated offenders

Proximity to motivated offenders captures situations in which respondents were directly exposed to fraud attempts or deceptive interactions during their online activities. Unlike exposure, which reflects routine use of online services, proximity focuses on direct contact with potential offenders, regardless of whether a concrete victimization outcome occurred.

The construct is operationalized using items from the Fraud module of the questionnaire and refers to events experienced in the previous six months. Proximity is derived from the following survey question:

“When using the Internet in the last 6 months, have you experienced any of the following situations?”

Respondents were asked to indicate whether they had encountered one or more specific situations related to online fraud or deceptive practices. The following items were included in the measurement of proximity to motivated offenders:

- **d1_pcan_1:** I have been asked for my user passwords or personal information
- **d1_pcan_2:** I have been offered a service or product that I did not request
- **d1_pcan_3:** I have been asked to click on a suspicious link or advertisement
- **d1_pcan_4:** I have received a job offer that I suspect may be fraudulent
- **d1_pcan_5:** I have been offered products from e-commerce websites that I suspect may be fake
- **d1_pcan_6:** I have accessed websites that were attempting to impersonate banks, retailers or public administrations
- **d1_pcan_7:** I have been signed up for services that I did not subscribe to

All items are coded as dichotomous indicators (0 = no / 1 = yes), reflecting whether the respondent experienced the corresponding situation during the reference period. Descriptive statistics for this construct, including the mean and standard deviation, are reported in the following table.

Table 4.3: L-RAT framework - Descriptive statistics of Proximity to motivated offender

Item	Mean	SD	Median	Min	Max	Skewness	Kurtosis	SE
Proximity (Average)	0.25	0.41	0	0	1	1.34	0.43	0.00
d1_pcan_1	0.19	0.39	0	0	1	1.56	0.45	0.00
d1_pcan_2	0.35	0.48	0	0	1	0.62	-1.62	0.01
d1_pcan_3	0.47	0.50	0	0	1	0.13	-1.98	0.01
d1_pcan_4	0.16	0.37	0	0	1	1.84	1.37	0.00
d1_pcan_5	0.33	0.47	0	0	1	0.71	-1.50	0.01
d1_pcan_6	0.11	0.31	0	0	1	2.53	4.39	0.00
d1_pcan_7	0.15	0.36	0	0	1	1.97	1.88	0.00

4.3.2.3. Target Suitability

Target suitability refers to the degree to which a user may appear attractive as a potential target for cyber offenders. This dimension captures individuals' willingness to disclose personal or sensitive information, which may facilitate the construction of detailed user profiles and increase vulnerability to cybercrime.

The construct is measured using items from the Opinion module of the questionnaire and focuses on users' confidence in providing personal information across different contexts. Online target suitability is derived from the following survey question:

"How confident do you feel about performing each of the following activities?"

The measurement is based on the item **g2_pcan**. Among all the items, the following five situations were selected::

- **g2_pcan_1**: Providing personal information online to register for a service (opening an email account, creating social media profiles, etc.)
- **g2_pcan_2**: Providing personal information at a private establishment (shop, bank, etc.) for a specific procedure
- **g2_pcan_3**: Providing personal information at a public entity (town hall, government agency, etc.) for a specific procedure
- **g2_pcan_4**: Providing personal information online on a public entity's website
- **g2_pcan_5**: Providing personal information via email or instant messaging

Responses were originally collected using categorical scales and were recoded so that higher values indicate higher target suitability, ranging from 1 = no confidence to 5 = much confidence. In the analytical sample, the mean value of the construct is 2.86, with a standard deviation of 0.97.

Table 4.4: L-RAT framework - Descriptive statistics of Target suitability

Item	Mean	SD	Median	Min	Max	Skewness	Kurtosis	SE
Suitability (Average)	3.14	0.97	3	1	5	-0.12	-0.27	0.01
g2_pcan_1	3.04	0.95	3	1	5	-0.03	-0.19	0.01
g2_pcan_2	3.10	0.95	3	1	5	-0.13	-0.22	0.01
g2_pcan_3	3.52	0.97	2	1	5	-0.34	-0.17	0.01
g2_pcan_4	3.36	0.98	3	1	5	-0.26	-0.28	0.01
g2_pcan_5	2.69	1.02	3	1	5	0.15	-0.49	0.01

Although this is an attitudinal construct, previous empirical work has shown that it is effective in predicting online risk, as it is positively associated with users' security

configurations and the visibility of their profiles on social networks (Herrero, Torres, Vivas, Hidalgo, et al., 2021).

4.3.2.4. Capable guardianship

Capable guardianship refers to users' perceived ability to protect themselves in digital environments and their awareness of online risks. This dimension captures attitudinal aspects related to cybersecurity, including perceived effort, economic and cognitive barriers to adopting security measures, and awareness of the consequences of online behaviour.

The construct is measured using items from the Opinion module of the questionnaire and is derived from the following survey question:

"Below we will show you various opinions on Internet security. Please tell us whether you agree or disagree with each one"

The measurement is based on the variable **g3_pcan**. Among all of them, the following five items were selected:

- **g3_pcan_1:** Using security tools and procedures to protect my computer/smartphone/tablet requires too much effort on my part
- **g3_pcan_3:** My lack of knowledge about security in new technologies limits my use of them
- **g3_pcan_5:** Security tools (antivirus, firewall, anti-spam, etc.) are too expensive for me
- **g3_pcan_7:** I would use more services via the Internet (banking, shopping, social networks) if I were taught how to do so safely
- **g3_pcan_9:** My online actions have consequences on cybersecurity

Respondents were asked to indicate their level of agreement with each statement using a five-point Likert scale, ranging from 1 = totally agree to 5 = totally disagree. Higher values indicate higher levels of capable guardianship, reflecting lower perceived barriers or high perceived competence in managing online security. In the analytical sample, the mean value of the construct is 2.45, with a standard deviation of 1.03.

Table 4.5: L-RAT framework - Descriptive statistics of Capable guardianship

Item	Mean	SD	Median	Min	Max	Skewness	Kurtosis	SE
Guardianship (Average)	2.45	1.02	2	1	5	0.49	-0.26	0.01
g3_pcan_1	2.31	1.01	2	1	5	0.62	-0.11	0.01
g3_pcan_3	2.62	1.07	2	1	5	0.37	-0.55	0.01
g3_pcan_5	2.38	1.00	2	1	5	0.53	-0.22	0.01
g3_pcan_7	2.41	1.02	2	1	5	0.48	-0.22	0.01
g3_pcan_9	2.53	1.02	2	1	5	0.45	-0.21	0.01

4.3.2.5. Modelling of the L-RAT framework as a higher-order construct

Each of the four L-RAT dimensions captures a distinct aspect of individuals' routine online activities. While each component contributes independently to cyber risk, the theoretical framework is intended to operate as a whole, affecting individuals through the combined effect of their online routines and security-related behaviours.

For this reason, the L-RAT framework is modelled as a Higher-Order Construct (HOC), integrating the four dimensions into a single analytical structure that captures their joint effect on cybervictimization risk. This approach is consistent with the theoretical literature on Lifestyle-Routine Activity Theory and its application to digital environments.

Following the strategy proposed by Sarstedt et al. (2019), the L-RAT construct is specified as a Type IV (formative-formative) HOC. The four lower-order components are not considered manifestations of a single latent trait, but rather distinct formative dimensions that jointly define the higher-order construct. Similarly, the L-RAT construct itself does not represent a homogeneous latent variable, but an analytical abstraction developed during the research design phase to delimit and operationalize the theoretical framework.

Given its formative nature, traditional reliability and internal consistency measures are not applicable to the higher-order construct. Instead, the evaluation of L-RAT follows established guidelines for formative measurement models. Each lower-order component is therefore assessed individually, in order to identify potential specification issues that could affect the validity of the higher-order construct and, consequently, the overall results of the study.

4.3.3. Gambling

Online gambling is included in the analysis as an observed behavioural variable, capturing respondents' direct participation in online betting and gambling activities. The variable reflects actual engagement in online gambling services, rather than attitudes or self-assessed gambling intensity.

The measure is derived from the internet services usage module of the questionnaire. Specifically, gambling behaviour is identified through the following survey question:

“From the following list of services offered on the Internet, indicate whether you have used them in the last six months from your computer/smartphone/tablet”

Within this list of services, respondents were asked to report their use of online betting or casino platforms. Individuals were classified as online gamblers if they responded “yes” to the item **a1_pcan_14**: “Online betting or casino services (e.g.: Bwin, Unibet, Miapuesta)”. Based on this item, the variable Gambling was constructed as a binary indicator, coded as:

- 0 = non-gambler, if the respondent did not report using online betting or gambling services
- 1 = gambler, if the respondent reported using online betting or gambling services

In the analytical sample, 777 respondents were identified as online gamblers. The descriptive statistics indicate a mean value of 0.11 and a standard deviation of 0.31, reflecting the relatively low prevalence of online gambling behaviour in the overall sample.

Table 4.6: Online gambling - Descriptive statistics

Item	Mean	SD	Median	Min	Max	Skewness	Kurtosis	SE
Gambling	0.11	0.31	0.00	0.00	1.00	2.49	4.21	0.00

Although gambling behaviour is measured using a single-item indicator, this operationalization is methodologically appropriate given the nature of the construct. Online gambling participation represents a concrete and unambiguous behavioural activity, for which single-item measures are considered acceptable and valid in empirical research.

Methodological literature on structural equation modelling supports the use of single-item measures when the construct is clearly defined, narrowly scoped, and directly observable, and when the objective is to capture factual behaviour rather than a latent psychological trait (Hair et al., 2021). In this context, the use of a dichotomous

indicator allows for a parsimonious and transparent representation of gambling involvement, while avoiding unnecessary measurement complexity.

Moreover, this specification facilitates the inclusion of gambling behaviour both as an observed predictor in the PLS-SEM model and as an input feature in the Machine Learning analyses, ensuring consistency across analytical approaches.

4.3.4. Smartphone Addiction

Smartphone addiction is included in the analysis as a reflective latent construct, capturing problematic and compulsive patterns of smartphone and tablet use. The construct is measured using eight items drawn from the Smartphone Addiction Symptoms Scale, which has been widely used in empirical research on problematic technology use (Bian & Leung, 2015; Young, 1998).

The selected items originate from two broader questionnaire variables, **i6aan** and **i6ban**, which assess smartphone and tablet usage habits. From these variables, the items most closely aligned with established screening instruments for Internet addiction were selected, following the strategy adopted in previous empirical studies (Herrero, Torres, Vivas, Arenas, et al., 2021).

Smartphone addiction is measured through the following survey question:

“Please specify to what extent the following statements reflect your use of your Smartphone/Tablet”

The construct is operationalized using the following items:

- **i6aan_4:** You have tried to hide how much time you spend on your smartphone/tablet from others
- **i6aan_5:** You have tried unsuccessfully to use your smartphone/tablet less
- **i6aan_9:** You have used your smartphone/tablet to feel better when you are feeling down
- **i6aan_10:** You spend more time on your smartphone/tablet than you had planned
- **i6ban_2:** You never get tired of being on your smartphone/tablet
- **i6ban_6:** You use your smartphone/tablet when you should really be doing other things, and this has caused you problems
- **i6ban_8:** When you lose connection, you are obsessed with whether you have missed calls
- **i6ban_9:** You feel anxious if you do not check your messages or if your smartphone/tablet is not turned on

All items are measured on a five-point Likert scale, ranging from 1 = never to 5 = most of the time, with higher values indicating higher levels of addictive behaviour.

In the analytical sample, the mean value of the scale is 2.46, with a standard deviation of 1.18. The internal consistency of the scale is high, with Cronbach's α equal to 0.905, indicating strong reliability. Unifactorial analysis was conducted and showed that all item loadings exceed the recommended threshold of 0.70, with the exception of item i6ban_2, which presents a loading of 0.628. Despite this lower value, the item was retained in the analysis due to its substantive contribution to the overall construct.

Table 4.7: *Smartphone addiction - Descriptive statistics*

Item	Mean	SD	Median	Min	Max	Skewness	Kurtosis	SE
SA (Average)	2.46	1.18	2	1	5	0.38	-0.74	0.01
i6aan_4	2.15	1.21	2	1	5	0.71	-0.61	0.01
i6aan_5	2.32	1.18	2	1	5	0.51	-0.70	0.01
i6aan_9	2.68	1.23	3	1	5	0.12	-0.99	0.01
i6aan_10	2.75	1.19	3	1	5	0.06	-0.91	0.01
i6ban_2	2.77	1.15	3	1	5	0.08	-0.75	0.01
i6ban_6	2.35	1.16	2	1	5	0.48	-0.70	0.01
i6ban_8	2.22	1.16	2	1	5	0.66	-0.49	0.01
i6ban_9	2.42	1.18	2	1	5	0.42	-0.76	0.01

4.3.5. Social Digital Pressure

Social Digital Pressure (SDP) captures the extent to which digital social environments generate expectations of availability, responsiveness, and digital competence. High levels of SDP indicate that individuals perceive stronger social pressure related to the use of digital devices, online applications, and social media platforms.

In the survey, SDP is measured using the question **i14_pcan**, which is based on the Social Digital Pressure Scale developed by Gui and Büchi (2021). This scale assesses perceived social expectations linked to communication practices and online presence. SDP is derived from the following survey question:

“In my daily life, the people I usually communicate with via social media and instant messaging (WhatsApp, Facebook, Instagram, Twitter, etc.) expect me to”

The construct is operationalized through the following three items:

- **i14_pcan_1:** Respond quickly to messages with my smartphone
- **i14_pcan_2:** Know how to use different applications
- **i14_pcan_3:** Stay active on social networks

Responses are collected on a five-point Likert scale, originally ranging from 1 = completely agree to 5 = completely disagree. Items were inverted so that higher values indicate higher levels of perceived social digital pressure.

The scale shows strong internal consistency, with Cronbach’s α equal to 0.850. Exploratory factor analysis confirmed a clear unifactorial structure, with all item loadings exceeding 0.82. Descriptive statistics indicate a mean value of 4.01 and a standard deviation of 1.10 in the analytical sample.

Table 4.8: Social digital pressure - Descriptive statistics

Item	Mean	SD	Median	Min	Max	Skewness	Kurtosis	SE
SDP (Average)	4.01	1.10	4	1	5	-0.85	-0.16	0.01
i14_pcan_1	4.07	1.08	4	1	5	-0.93	-0.01	0.01
i14_pcan_2	4.09	1.03	4	1	5	-0.94	0.18	0.01
i14_pcan_3	3.88	1.19	4	1	5	-0.68	-0.64	0.01

4.3.6. Social Support

Social support refers to the perception, and in some cases the actual experience, of being cared for and embedded in a network of supportive relationships, such as family members, friends, neighbours, or other close contacts. Lower levels of social support indicate a weaker perception of belonging to a supportive and caring social environment.

In this study, social support is measured as a reflective latent construct using Lin’s Strong-Tie Support Scale (Lin et al., 1981), which focuses on perceived support derived from close and trusted relationships. The scale is designed for large-scale surveys and has been shown to provide a reliable estimation of perceived social support due to its simplicity and stable psychometric properties (Herrero, Torres, Vivas, Arenas, et al., 2021).

Social support is operationalized using the survey question **i9_pcan**, formulated as follows:

“Please specify how often the following problems have concerned you in the last 6 months”

The construct is measured through the following three items:

- **i9_pcan_1:** Not having a partner
- **i9_pcan_2:** Not seeing family and friends
- **i9_pcan_3:** Not having close friends

Responses are collected on a five-point Likert scale, originally ranging from 1 = never to 5 = most of the time. Items were inverted so that higher values indicate higher levels of perceived social support. The scale shows acceptable internal consistency, with Cronbach’s α equal to 0.673. Descriptive statistics indicate a mean value of 3.48 and a standard deviation of 1.22 in the analytical sample. All items exhibit factor loadings above 0.70, with the exception of item i9pcan_2, which presents a loading of 0.680. Despite this lower value, the item was retained in the analysis due to its contribution to the overall construct.

Table 4.9: Social support - Descriptive statistics

Item	Mean	SD	Median	Min	Max	Skewness	Kurtosis	SE
SS (Average)	3.48	1.22	4	1	5	-0.36	-0.58	0.01
i9_pcan_1	4.00	1.24	5	1	5	-1.01	-0.13	0.01
i9_pcan_2	2.87	1.19	3	1	5	0.30	-0.74	0.01
i9_pcan_3	3.56	1.22	4	1	5	-0.38	-0.87	0.01

4.3.7. Sociodemographic variables

Several sociodemographic variables were included in the analysis as control variables, in order to account for potential confounding effects and reduce unexplained variability in the empirical models. Although the questionnaire collects a wide range of background information, the present study focuses on age, gender, and educational level, as these dimensions represent well-established and stable indicators in behavioural research and are commonly used in large-scale survey studies.

The choice of these variables is also motivated by the availability of official and disaggregated population statistics, which allows for a direct comparison between the sample distribution and national demographic figures. This comparison supports an assessment of the representativeness of the analytical sample.

4.3.7.1. Age

Age was originally collected in the survey as a continuous numerical variable and was subsequently recoded into five age groups for operational purposes:

1. **15-24 years:** 10.7%
2. **25-34 years:** 19.0%
3. **35-44 years:** 27.2%
4. **45-54 years:** 23.5%
5. **55 years and older:** 19.6%

Overall, the age distribution of the sample broadly resembles that of the Spanish population, although some discrepancies emerge when compared with official statistics provided by the Spanish National Institute for Statistics, updated as of 1 July 2021 (Instituto Nacional de Estadística, 2021). Individuals aged 25-34, 35-44, and 45-54 are slightly overrepresented, while the youngest (15-24) and oldest (55+) groups are more closely aligned with national figures.

These differences are partly attributable to the characteristics of the target population. National population statistics include children and very elderly individuals, whereas the present sample is restricted to respondents who own and use digital devices. Consequently, the observed deviations are more likely related to differences in technology adoption across age groups rather than to systematic sampling bias.

4.3.7.2. Gender

Gender is coded as a binary variable, with 1 = male and 2 = female. In the analytical sample, 46.7% of respondents identified as men and 53.3% as women. This distribution is broadly consistent with national demographic figures.

4.3.7.3. Educational level

Educational background was measured by asking respondents to report the highest level of education attained. Original response categories were recoded into four educational groups, reflecting increasing levels of educational attainment:

1. **Primary education:** 2.0%
2. **Secondary education (up to high school):** 49.2%
3. **Higher education (university degree or higher):** 48.7%
4. **Other educational categories**

When compared with national statistics, the sample shows a lower proportion of individuals with primary education (2.0% versus 10.8% nationally), a similar proportion of respondents with secondary education (49.2% versus 51.3%), and an

overrepresentation of individuals with higher education (48.7% in the sample compared to 32.2% in the national population).

As with age, these differences are primarily attributed to unequal patterns of digital technology adoption, which tend to be higher among individuals with higher educational attainment.

Table 4.10: *Sociodemographic variables - Descriptive statistics*

Item	Mean	SD	Median	Min	Max	Skewness	Kurtosis	SE
cont_edad_int	3.22	1.26	3	1	5	-0.16	-0.98	0.01
cont_sexo	1.53	0.5	2	1	2	-0.13	-1.98	0.01
cont_estudio	2.47	0.54	2	1	4	-0.24	-1.08	0.01

4.4. Structural Modelling Strategy: PLS-SEM

The primary analytical approach adopted in this study is Partial Least Squares Structural Equation Modelling (PLS-SEM). PLS-SEM is a variance-based method designed to estimate complex models involving multiple latent variables and hierarchical measurement structures, and is particularly suitable for predictive-oriented research and models combining reflective and formative constructs (Hair et al., 2021).

Latent variables are operationalized according to the specifications presented in Section 4.3, using indicators selected from the survey questionnaire. The objective of the PLS-SEM model is to assess which factors contribute to cybervictimization risk, to evaluate the relative magnitude of these effects, and to determine the statistical significance of each component within the overall structural model.

The modelling strategy incorporates a hierarchical specification to represent the broader Lifestyle-Routine Activity Theory (L-RAT) framework. In line with the theoretical rationale, the L-RAT construct is specified as a formative-formative higher-order construct (HOC), allowing the joint contribution of its components to be assessed within a unified analytical structure. This approach reduces model complexity while improving convergence and interpretability of the results.

All statistical analyses are conducted using the open-source software RStudio version 2024.12.1 Build 563.

4.4.1. Operative Strategy: First-Order Model specification

The analytical procedure begins with the specification of the measurement model, in which survey items are assigned to their corresponding latent constructs. At

this stage, a first-order (“base”) model is estimated, including only first-order constructs. This step is necessary to ensure a correct measurement specification before proceeding to the estimation of the higher-order L-RAT construct.

Constructs are specified as reflective or formative based on their theoretical foundations and measurement logic. Smartphone Addiction (SA), Social Digital Pressure (SDP), and Social Support (SS) are modelled as reflective constructs (*Mode A*), as they represent latent traits measured through multiple correlated indicators. In contrast, the four L-RAT components (Exposure, Proximity, Target Suitability, and Capable Guardianship) are modelled as formative constructs (*Mode B*), as each dimension is defined by distinct behavioural indicators that do not necessarily covary and jointly form the construct.

Once the measurement model is defined, the first-order structural model is specified. Structural paths are defined a priori, reflecting the theoretical relationships discussed in Chapter 2 and the methodological design presented in Chapter 3. Within the structural model, constructs are classified as exogenous or endogenous. Social Digital Pressure and Social Support are treated as exogenous constructs, while all other variables are considered endogenous, as they are both predicted by earlier constructs and act as predictors of subsequent outcomes, including cybervictimization risk.

The first-order PLS-SEM model is estimated using the *semnir* package in R, which relies on ordinary least squares regressions to iteratively estimate outer weights, loadings, and structural path coefficients. To assess the stability and statistical significance of the estimated parameters, a non-parametric bootstrapping procedure with 10,000 resamples is applied. Bootstrapping enables the computation of standard errors, confidence intervals, and bias-corrected p-values for all estimated effects.

4.4.2. Operative Strategy: Higher-Order Model Specification

The estimation of the first-order model represents the foundation for the construction of the L-RAT higher-order construct. Following the disjoint two-stage approach, latent variable scores for the four lower-order components (Exposure, Proximity, Target Suitability, and Capable Guardianship) are extracted from the first-order model and included in the dataset as pseudo-indicators.

In the second stage, the measurement model is re-specified to include the formative L-RAT construct, and the structural model is adjusted accordingly. The full model is then re-estimated using the same PLS-SEM procedure and bootstrapped again with 10,000 resamples.

Given the formative nature of the higher-order construct, traditional reliability and internal consistency measures are not applicable. Instead, the evaluation follows established guidelines for formative measurement models, focusing on indicator relevance, collinearity diagnostics, and the contribution of each lower-order component to the higher-order construct (Hair et al., 2021; Sarstedt et al., 2019).

This two-stage hierarchical modelling strategy allows the study to capture the overall effect of L-RAT on cybervictimization risk while preserving the conceptual distinctiveness of its underlying dimensions.

4.5. Predictive Modelling Strategy: Machine Learning

In addition to the PLS-SEM analysis, this study adopts a Machine Learning (ML) approach to examine cybervictimization risk from a purely predictive perspective. Unlike structural equation modelling, machine learning techniques do not rely on the explicit specification of latent variables, causal paths, or higher-order constructs. Instead, all observed variables are treated as direct predictors of the outcome, and relationships are inferred directly from the data.

The two approaches serve distinct but complementary objectives. PLS-SEM is employed to test theoretically grounded hypotheses and to estimate the magnitude and significance of predefined structural relationships. In contrast, the Machine Learning analysis aims to explore predictive relevance and to assess the relative contribution of each variable to cybervictimization risk without imposing a priori structural constraints. In this sense, ML is not used for causal inference, but as a data-driven validation tool that complements the theory-based SEM results.

Following the methodological framework outlined in Section 3.3.2, the selected algorithm for the Machine Learning analysis is Random Forest, due to its robustness, ability to handle non-linear relationships, and effectiveness in capturing complex interactions among predictors. Random Forest has been widely used in behavioural and social science research when the primary goal is prediction rather than explanation, and when variables may exhibit heterogeneous and non-additive effects.

The main objective of this section is therefore to compare predictive performance and variable importance patterns derived from the Random Forest model with the results obtained through PLS-SEM.

4.5.1. Operative strategy: Variable definition and feature construction

The Machine Learning analysis uses the same set of variables that define the final PLS-SEM model, ensuring full comparability across approaches. However, due to the absence of latent variable modelling in ML, constructs are operationalized as observed features, following the strategy described in the methodology chapter.

Reflective constructs (i.e., Smartphone Addiction, Social Digital Pressure, and Social Support) are operationalized as single composite variables, computed as the arithmetic mean of their respective items (*pomp_mean*). The four components of the L-RAT framework (i.e., Exposure, Proximity, Target Suitability, and Capable Guardianship) which are specified as formative constructs in the SEM model, are

operationalized as summed indices, reflecting the aggregation logic used in their original measurement (*pomp_sum*).

To ensure comparability across variables and to avoid scale-related distortions, all predictors above are rescaled using a Percent of Maximum Possible (POMP) transformation, mapping values onto a 0-100 range, as recommended in the methodological literature for machine learning applications in social research (Boehmke & Greenwell, 2019). The gambling variable, originally coded as a binary indicator (0 = non-gambler, 1 = gambler), was rescaled to a 0-100 range to ensure comparability with the other predictors used in the model. Given its single-item and behavioural nature, the variable was not aggregated using POMP mean or sum procedures, but linearly transformed to preserve its dichotomous interpretation.

Sociodemographic control variables (i.e., gender, age, and educational level) are encoded using dummy variables, applying a one-hot encoding strategy where appropriate. The dependent variable, cybervictimization risk, is also rescaled to the 0-100 range, allowing consistency in model estimation and interpretation of prediction errors.

Table 4.11: Machine Learning dataset - Descriptive statistics

Item	Mean	SD	Median	Min	Max	Skewness	Kurtosis	SE
pc1a_Riesgo	6.22	13.38	0.00	0	100	2.90	10.61	0.16
Exposure	46.52	27.88	50.00	0	100	0.05	-0.68	0.33
Proximity	25.21	23.81	14.29	0	100	0.78	-0.04	0.28
Suitability	53.55	18.85	50.00	0	100	-0.20	0.20	0.22
Guardianship	36.26	16.87	35.00	0	100	0.28	0.35	0.20
Gambling	11.00	31.29	0.00	0	100	2.49	4.21	0.37
SA	61.95	23.60	66.67	0	100	-0.32	-0.32	0.28
SDP	75.37	24.05	75.00	0	100	-0.72	-0.10	0.29
SS	36.42	22.92	34.38	0	100	0.41	-0.49	0.27

4.5.2. Operative strategy: Training, validation, and model estimation

The dataset is split into training and test subsets, following the procedure described in Chapter 3. A proportion of 80% of the observations is used for model training, while the remaining 20% is reserved for out-of-sample evaluation. To reduce

potential distortions due to imbalanced outcome distributions, the split is performed using stratification on the binary cybervictimization outcome (*Victim* vs. *NonVictim*).

Model training and validation are conducted using the procedure described by Boehmke & Greenwell (2019). The *tidymodels* package is used since it provides a unified and reproducible framework for resampling, tuning, and evaluation. Cross-validation is implemented using a k-fold strategy with $k = 10$, ensuring stable estimation of predictive performance and limiting overfitting.

Several Random Forest hyperparameters are tuned during the training phase, following a grid-search strategy defined a priori. Other parameters, instead, are held fixed to ensure methodological consistency with the classical Random Forest implementation. In particular, the sampling fraction was fixed at 0.632, corresponding to the standard bootstrap proportion traditionally used in Random Forest algorithms (Breiman, 2001). Additionally, sampling was performed with replacement (*replace = TRUE*), in accordance with the original bootstrap aggregation procedure. While the remaining hyperparameters were optimized through cross-validation, these settings were kept constant to preserve the theoretical properties of bootstrap aggregation underlying the Random Forest algorithm.

Model selection is based on classification-oriented performance metrics, in line with the evaluation criteria defined in the methodology section. Once the optimal model configuration is identified, the final model is fitted on the training data and subsequently evaluated on the test set.

The “best” Random Forest model identified is composed of the following parameters:

- number of trees: *n_trees* = 1844
- number of variables randomly sampled at each split: *mtry* = 5
- maximum tree depth: *max_depth* = 6
- minimum number of observations in a terminal node: *min_n* = 25

4.5.3. Operative strategy: Model evaluation and interpretation

After training, the Random Forest model is evaluated on the test dataset to assess out-of-sample predictive performance. Model fit is examined using multiple error metrics, and predictions obtained from the training and test samples are compared to evaluate model stability.

In addition to predictive accuracy, the analysis focuses on variable importance measures, which provide information on the relative contribution of each predictor to the model’s performance. These importance rankings are used to identify the most influential factors associated with cybervictimization risk and to compare data-driven relevance with the theoretical structure estimated through PLS-SEM.

Finally, once the model has been validated, it is re-estimated on the full dataset to obtain stable importance measures and final predictions. The results of the Machine Learning analysis are presented and discussed in conjunction with the SEM findings, allowing for an integrated interpretation of explanatory and predictive evidence.

5 Results and discussion

This chapter presents and discusses the empirical results of the study. First, the results of the PLS-SEM model are reported, including the assessment of the measurement model, the structural model, and the mediation analysis. Second, the findings of the Machine Learning approach are presented, focusing on predictive performance and variable importance. The chapter concludes with an integrated discussion of the results, linking empirical evidence to the theoretical framework and the research hypotheses formulated in Chapter 2.

5.1. PLS-SEM results

The evaluation of PLS-SEM results follows the two-step approach described in Chapter 3, beginning with the assessment of the measurement model and proceeding with the analysis of the structural relationships and mediation effects. The objective is to test the theoretical hypotheses and to examine the explanatory mechanisms underlying cybervictimization risk.

5.1.1. Reflective measurement model

The assessment of the reflective measurement model concerns exclusively the constructs specified in Mode A, namely Smartphone Addiction (SA), Social Digital Pressure (SDP) and Social Support (SS). In line with the procedure described in Chapter 3, the evaluation follows four sequential steps: indicators reliability, internal consistency reliability, convergent validity and discriminant validity (Hair et al., 2019, 2021).

Indicators reliability is assessed through the examination of outer loadings and their statistical significance obtained via bootstrapping. As reported in Table 5.1, the vast majority of indicators load above the recommended threshold of 0.70. Only two items display loadings slightly below 0.70 but remain above 0.60, which is considered acceptable in predictive-oriented PLS-SEM applications when theoretical justification is strong. All outer loadings are statistically significant based on bootstrapped confidence intervals, confirming that each indicator contributes meaningfully to its respective latent construct. No problematic cross-loading patterns emerge, supporting the unidimensionality of SA, SDP and SS.

Internal consistency reliability is evaluated using Cronbach's Alpha, Composite Reliability (rhoC), and rhoA. As shown in the table, all three constructs present values above the recommended threshold of 0.70 for all reliability indices. CR and rhoA values confirm that the indicators jointly provide a consistent representation of their respective latent variables. Importantly, none of the constructs exhibits excessively high reliability values (i.e., above 0.95), thus excluding redundancy among indicators. These results confirm satisfactory internal consistency for SA, SDP and SS.

Table 5.1: Reflective measurement model results

Construct /indicator	Total Sample (N = 7,064)			
	Outer loadings (Alpha)	rhoA	rhoC	AVE
SA	(0.905)	0.912	0.924	0.604
i6aan_4	0.793			
i6aan_5	0.792			
i6aan_9	0.770			
i6aan_10	0.783			
i6ban_2	0.628			
i6ban_6	0.836			
i6ban_8	0.764			
i6ban_9	0.829			
SDP	(0.850)	0.921	0.906	0.764
i14pcan_1	0.879			
i14pcan_2	0.915			
i14pcan_3	0.823			
SS	(0.673)	0.700	0.820	0.605
i9pcan_1	0.784			
i9pcan_2	0.680			
i9pcan_3	0.859			

Convergent validity is assessed through the Average Variance Extracted (AVE), which should exceed 0.50 to indicate that the construct explains more than half of the variance of its indicators. All reflective constructs report AVE values above the 0.50 threshold, confirming adequate convergent validity. This indicates that the indicators associated with each latent variable share a high proportion of common variance and effectively capture the intended theoretical dimension.

Discriminant validity is evaluated using two complementary criteria: the Fornell-Larcker criterion and the Heterotrait-Monotrait ratio (HTMT). First, according to the Fornell-Larcker criterion, the square root of the AVE of each construct exceeds its correlations with other constructs, confirming that each latent variable shares more variance with its own indicators than with other constructs in the model. Second, the HTMT ratios remain below the conservative threshold of 0.85 (and well below 0.90 in all cases), and the corresponding bootstrap confidence intervals do not include the value 1. This provides further evidence that SA, SDP and SS are empirically distinct constructs. Results are provided in Table 5.2.

Table 5.2: *Reflective measurement model results: discriminant validity*

Total Sample (N = 7,064)							
Fornell-Larcker criterion*				HTMT**			
	SA	SDP	SS		SA	SDP	SS
SA	0.777			SA			
SDP	-0.092	0.874		SDP	0.098 [0.073 ; 0.125]		
SS	-0.405	0.083	0.778	SS	0.508 [0.479 ; 0.536]	0.103 [0.072 ; 0.134]	

* the **emboldened** numbers on the diagonal represent the square root of the AVE for each construct; the remaining numbers are the correlations between the reflective constructs from the latent scores.

** In parentheses, the 95% CI.

Overall, the reflective measurement model demonstrates satisfactory psychometric properties. Indicator reliability, internal consistency reliability, convergent validity and discriminant validity are all supported for SA, SDP and SS. These results ensure that the reflective constructs are measured in a statistically robust and theoretically coherent manner, thereby providing a reliable basis for the subsequent evaluation of the formative higher-order construct (L-RAT) and for the structural model assessment.

5.1.2. Formative measurement model

The formative measurement assessment concerns the constructs modelled in Mode B within the L-RAT framework. In line with the theoretical specification presented in Chapter 3, the four dimensions Exposure, Proximity, Target Suitability, and Capable Guardianship are treated as formative components, as they represent distinct and non-interchangeable elements that jointly define the theoretical construct.

The evaluation of formative constructs differs from that of reflective ones and focuses primarily on (i) collinearity among indicators and (ii) the magnitude and statistical significance of indicator weights (Hair et al., 2019, 2021).

As reported in Table 5.3, all Outer VIF values remain well below the conservative threshold of 3, excluding problematic multicollinearity. This result indicates that the formative indicators contribute unique information and that the estimation of weights is not distorted by redundancy among predictors. The analysis of indicator weights shows that most formative indicators display statistically significant contributions, as their bias-corrected 95% confidence intervals do not include zero. This confirms that the majority of indicators contribute meaningfully to the construction of their respective formative dimensions. However, two indicators present non-significant weights. In accordance with PLS-SEM guidelines for formative constructs, these indicators are not automatically removed. Instead, their outer loadings are examined to assess whether they still share substantial variance with the construct.

For both d1_pcan_2 and g2_pcan_3, the reported loadings are statistically significant, as their BC 95% confidence intervals do not include zero. This indicates that, although their relative contribution (weight) is limited, they maintain a meaningful absolute contribution to the construct. Given their theoretical relevance within the L-RAT framework and the absence of collinearity issues, both indicators are retained in the final model.

Table 5.3: Formative measurement model results

Dimension / indicator	Total Sample (N = 7,064)				
	Outer VIF	Weight	Weight's BC CI 95%	Loading	Loading's BC CI 95%
Exposure	1.038	0.359 ***	[0.309 ; 0.407]		
a1_pcan_6	1.06	0.286 ***	[0.214 ; 0.357]		
a1_pcan_7	1.06	-0.146 ***	[-0.227 ; -0.068]		
a1_pcan_8	1.10	0.670 ***	[0.604 ; 0.733]		
a1_pcan_13	1.11	0.524 ***	[0.450 ; 0.596]		
Guardian	1.045	-0.359 ***	[-0.403 ; -0.315]		
g3_pcan_1	1.36	-0.156 ***	[-0.246 ; -0.065]		
g3_pcan_3	1.30	0.527 ***	[0.442 ; 0.610]		
g3_pcan_5	1.36	0.156 ***	[0.067 ; 0.245]		
g3_pcan_7	1.31	0.325 ***	[0.236 ; 0.410]		
g3_pcan_9	1.35	0.467 ***	[0.395 ; 0.538]		
Proximity	1.032	0.606 ***	[0.563 ; 0.649]		
d1_pcan_1	1.26	0.257 ***	[0.186 ; 0.327]		
d1_pcan_2	1.22	0.020 n.s.	[-0.048 ; 0.087]	0.342 ***	[0.283 ; 0.400]
d1_pcan_3	1.21	-0.113 ***	[-0.179 ; -0.046]		
d1_pcan_4	1.27	0.327 ***	[0.257 ; 0.397]		
d1_pcan_5	1.21	0.198 ***	[0.126 ; 0.270]		
d1_pcan_6	1.27	0.472 ***	[0.397 ; 0.542]		
d1_pcan_7	1.28	0.462 ***	[0.395 ; 0.527]		
TargetSuit	1.047	0.393 ***	[0.349 ; 0.437]		
g2_pcan_1	2.37	0.344 ***	[0.234 ; 0.455]		
g2_pcan_2	2.33	0.310 ***	[0.193 ; 0.421]		
g2_pcan_3	1.86	-0.128 n.s.	[-0.271 ; 0.014]	0.151 ***	[0.062 ; 0.240]
g2_pcan_4	1.89	-0.357 ***	[-0.492 ; -0.217]		
g2_pcan_5	2.36	0.727 ***	[0.644 ; 0.807]		

Note: The following convention values have been adopted for significance levels:

- t-values below 1.96 are considered not significant at the 5% level ($p > 0.05$)
- values between 1.96 and 2.58 are marked with * ($p < 0.05$)
- values between 2.58 and 3.29 are marked with ** ($p < 0.01$)
- values above 3.29 are marked with *** ($p < 0.001$)

Overall, the formative measurement results confirm that the L-RAT dimensions are specified in a statistically coherent and theoretically consistent manner. The absence of multicollinearity and the predominance of significant weights, together with the justified retention of the two indicators with non-significant weights, provide sufficient support for proceeding to the evaluation of the structural model.

5.1.3. Structural model

Once the measurement model has been validated, the next step consists in evaluating the structural model. In line with the procedure described in Chapter 3, the assessment focuses on (i) collinearity among predictor constructs, (ii) significance and magnitude of path coefficients, (iii) explanatory power (R^2), and (iv) effect sizes (f^2) (Hair et al., 2019, 2021). The structural relationships are evaluated both in the first-order specification and in the second-order specification, where the four L-RAT dimensions are modelled jointly through the higher-order construct.

Collinearity among predictor constructs is assessed through Inner VIF values computed on latent variable scores. In both model specifications, Inner VIF values remain below the conservative threshold of 3, indicating the absence of problematic multicollinearity. This result ensures that path coefficient estimates are stable and not inflated due to redundancy among predictors. For the path from SA to Gambling, the Inner VIF is reported as NA. This is expected, as Gambling is predicted by a single exogenous construct. Since collinearity assessment requires at least two predictors for a given endogenous construct, the VIF statistic is not applicable in this case.

The estimated path coefficients and their bias-corrected 95% bootstrap confidence intervals are reported in Table 5.4 and Table 5.4

In the first-order specification, the four L-RAT dimensions (Exposure, Proximity, Target Suitability, and Capable Guardianship), together with Smartphone Addiction (SA), Social Digital Pressure (SDP), and Gambling, are simultaneously included as predictors of cybervictimization risk. The results indicate that:

- Smartphone Addiction and Gambling show positive and statistically significant direct associations with cybervictimization risk
- among the L-RAT dimensions, Proximity, Exposure, and Target Suitability are positively and significantly related to the outcome variable, whereas Capable Guardianship displays a negative and statistically significant relationship
- the path from Gambling to Guardianship does not reach statistical significance, while all other structural paths within the L-RAT components are significant.

Statistical significance is assessed through bootstrapping procedures, and relevance is determined based on the bias-corrected 95% confidence intervals reported in the table.

Table 5.4: Structural Model results, first-order model

Endogenous construct	Relationship		Total Sample (N = 7,064)		
			Inner VIF	Path coef β	BC CI 95%
Smartphone Addiction	SS	→ SA	1.007	-0.400	[-0.422 ; -0.377]
	SDP	→ SA	1.007	-0.058	[-0.079 ; -0.037]
Gambling	SA	→ Gambling	NA	0.137	[0.112 ; 0.162]
Exposure	SA	→ Exposure	1.019	0.102	[0.076 ; 0.126]
	Gambling	→ Exposure	1.019	0.243	[0.214 ; 0.269]
Proximity	SA	→ Proximity	1.019	0.180	[0.156 ; 0.203]
	Gambling	→ Proximity	1.019	0.112	[0.083 ; 0.140]
TargetSuit	SA	→ TargetSuit	1.019	0.252	[0.226 ; 0.274]
	Gambling	→ TargetSuit	1.019	0.076	[0.050 ; 0.101]
Guardian	SA	→ Guardian	1.019	-0.266	[-0.288 ; -0.242]
	Gambling	→ Guardian	1.019	-0.022	[-0.045 ; 0.001]
Target Variable pc1a_Riesgo	SA	→ pc1a_Riesgo	1.296	0.168	[0.144 ; 0.193]
	Gambling	→ pc1a_Riesgo	1.110	0.065	[0.038 ; 0.094]
	Exposure	→ pc1a_Riesgo	1.130	0.087	[0.061 ; 0.113]
	Proximity	→ pc1a_Riesgo	1.084	0.292	[0.262 ; 0.319]
	TargetSuit	→ pc1a_Riesgo	1.100	0.050	[0.026 ; 0.072]
	Guardian	→ pc1a_Riesgo	1.096	-0.067	[-0.088 ; -0.046]
	Age	→ pc1a_Riesgo	1.136	-0.063	[-0.087 ; -0.040]
	Gender	→ pc1a_Riesgo	1.062	-0.011	[-0.033 ; 0.010]
	Education	→ pc1a_Riesgo	1.008	-0.006	[-0.027 ; 0.015]

In the second-order specification, the four L-RAT dimensions are integrated into a single higher-order construct (L-RAT), which is then used as predictor of cybervictimization risk. As reported in Table 5.5, the higher-order L-RAT construct exhibits a positive and statistically significant relationship with cybervictimization risk. Smartphone Addiction also remains statistically significant in this specification.

Table 5.5: Structural Model results, second-order model

Endogenous construct	Relationship		Total Sample (N = 7,064)		
			Inner VIF	Path coef β	BC CI 95%
Smartphone Addiction	SS	→ SA	1.007	-0.400	[-0.422 ; -0.378]
	SDP	→ SA	1.007	-0.058	[-0.079 ; -0.037]
Gambling	SA	→ Gambling	NA	0.137	[0.112 ; 0.162]
LRAT	SA	→ LRAT	1.019	0.340	[0.315 ; 0.362]
	Gambling	→ LRAT	1.019	0.193	[0.164 ; 0.221]
Target Variable pc1a_Riesgo	SA	→ pc1a_Riesgo	1.263	0.151	[0.126 ; 0.175]
	Gambling	→ pc1a_Riesgo	1.084	0.062	[0.034 ; 0.091]
	LRAT	→ pc1a_Riesgo	1.215	0.321	[0.291 ; 0.351]
	Age	→ pc1a_Riesgo	1.127	-0.062	[-0.086 ; -0.038]
	Gender	→ pc1a_Riesgo	1.048	-0.009	[-0.031 ; 0.013]
	Education	→ pc1a_Riesgo	1.004	-0.002	[-0.023 ; 0.019]

Among the control variables, only Age shows a statistically significant relationship with cybervictimization risk. In contrast, Gender and Education do not exhibit statistically significant path coefficients in both specifications.

The explanatory power of the model is evaluated through the coefficient of determination (R^2), which represents the proportion of variance explained in each endogenous construct. The results are reported in Table 5.6 and Table 5.7.

Bootstrapping procedures were applied to compute confidence intervals for R^2 , as described in Chapter 3. The resulting values confirm the stability of the explained variance estimates across specifications. Overall, both model configurations display consistent explanatory performance, with no relevant loss of predictive capacity when moving from the first-order to the second-order specification.

Table 5.6: *Explanatory power (R^2) results, first-order model*

Construct	Total sample (N = 7,064)	
	R^2 value	BC CI 95%
SA	0.168	[0.158 ; 0.176]
Gambling	0.019	[0.015 ; 0.022]
Exposure	0.076	[0.069 ; 0.083]
Proximity	0.050	[0.045 ; 0.056]
TargetSuit	0.074	[0.068 ; 0.081]
Guardian	0.073	[0.066 ; 0.078]
Target Variable pc1a_Riesgo	0.212	[0.200 ; 0.222]

Table 5.7: *Explanatory power (R^2) results, second-order model*

Construct	Total sample (N = 7,064)	
	R^2 value	BC CI 95%
SA	0.168	[0.159 ; 0.176]
Gambling	0.019	[0.015 ; 0.022]
LRAT	0.171	[0.161 ; 0.179]
Target Variable pc1a_Riesgo	0.194	[0.183 ; 0.205]

To assess the relative contribution of each predictor to the endogenous construct, effect sizes (f^2) are examined and interpreted according to standard thresholds (0.02 small, 0.15 medium, 0.35 large). Results are reported in Table 5.8 and Table 5.9.

In the first-order model, the f^2 indicate that certain predictors contribute more substantially to the explained variance of cybervictimization risk, while others show smaller incremental contributions. The removal of predictors with very small f^2 values would not substantially reduce the model's explanatory power. In the second-order model, the aggregated L-RAT construct exhibits an effect size consistent with its role as a central explanatory component of the framework. Smartphone Addiction and Gambling also show meaningful contributions, whereas the remaining predictors display smaller incremental effects.

The comparison between the two specifications confirms that the higher-order modelling strategy preserves the substantive contribution of the routine activity dimensions while providing a more parsimonious structural representation.

Table 5.8: Effect size (f^2) results, first-order model

Endogenous construct	Relationship			Total Sample (N = 7,064)			
				R ² included	R ² excluded	f ² value	Effect size
Smartphone Addiction	SS	→	SA	0.168	0.005	1.950e-01	medium
	SDP	→	SA	0.168	0.162	6.404e-03	negligible
Gambling	SA	→	Gambling	0.019	0.000	1.917e-02	negligible
Exposure	SA	→	Exposure	0.076	0.062	1.488e-02	negligible
	Gambling	→	Exposure	0.076	0.014	6.752e-02	small
Proximity	SA	→	Proximity	0.050	0.015	3.703e-02	small
	Gambling	→	Proximity	0.050	0.035	1.606e-02	negligible
TargetSuit	SA	→	TargetSuit	0.074	0.008	7.116e-02	small
	Gambling	→	TargetSuit	0.074	0.066	9.058e-03	negligible
Guardian	SA	→	Guardian	0.073	0.001	7.694e-02	small
	Gambling	→	Guardian	0.073	0.071	1.356e-03	negligible
Target Variable pc1a_Riesgo	SA	→	pc1a_Riesgo	0.212	0.162	6.303e-02	small
	Gambling	→	pc1a_Riesgo	0.212	0.201	1.346e-02	negligible
	Exposure	→	pc1a_Riesgo	0.212	0.195	2.194e-02	small
	Proximity	→	pc1a_Riesgo	0.212	0.106	1.349e-01	small
	TargetSuit	→	pc1a_Riesgo	0.212	0.205	8.282e-03	negligible
	Guardian	→	pc1a_Riesgo	0.212	0.201	1.377e-02	negligible
	Age	→	pc1a_Riesgo	0.212	0.201	1.328e-02	negligible
	Gender	→	pc1a_Riesgo	0.212	0.212	9.980e-05	negligible
	Education	→	pc1a_Riesgo	0.212	0.212	-9.284e-05	negligible

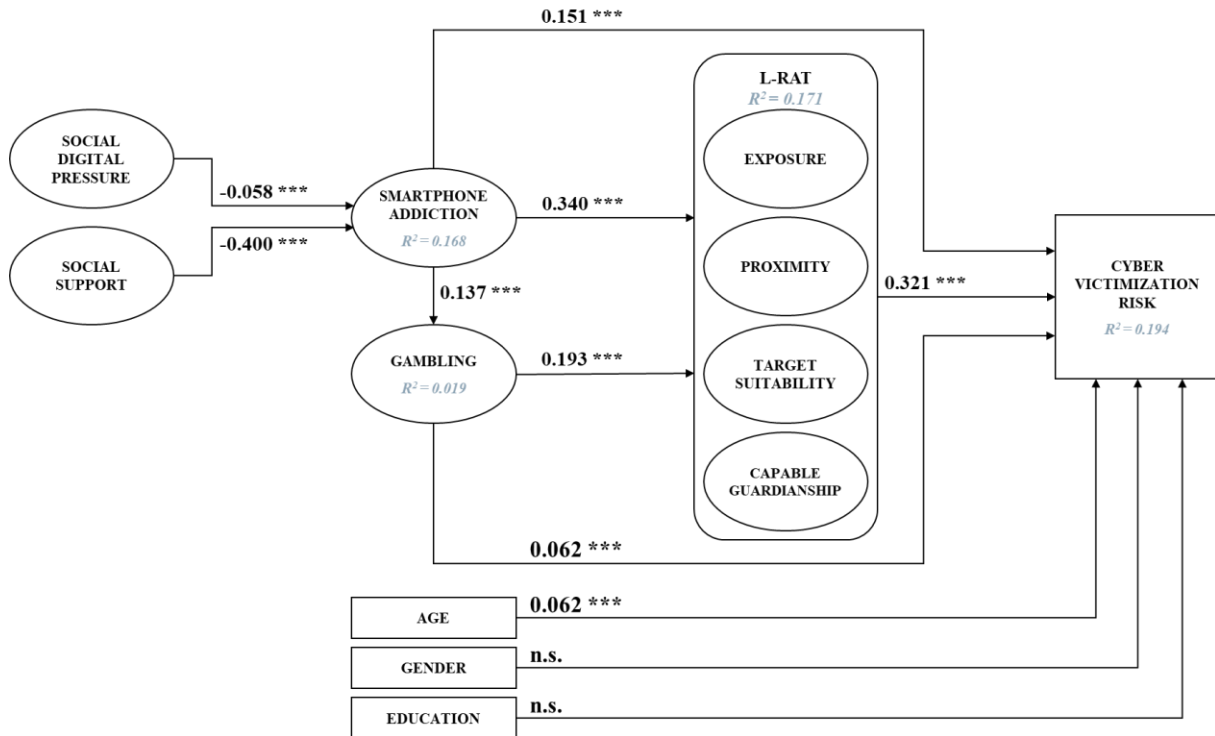
Table 5.9: Effect size (f^2) results, second-order model

Endogenous construct	Relationship			Total Sample (N = 7,064)			
				R ² included	R ² excluded	f ² value	Effect size
Smartphone Addiction	SS	→	SA	0.168	0.005	1.950e-01	medium
	SDP	→	SA	0.168	0.162	6.390e-03	negligible
Gambling	SA	→	Gambling	0.019	0.000	1.914e-02	negligible
LRAT	SA	→	LRAT	0.171	0.046	1.502e-01	medium
	Gambling	→	LRAT	0.171	0.125	5.567e-02	small
Target Variable pc1a_Riesgo	SA	→	pc1a_Riesgo	0.194	0.150	5.514e-02	small
	Gambling	→	pc1a_Riesgo	0.194	0.184	1.261e-02	negligible
	LRAT	→	pc1a_Riesgo	0.194	0.065	1.604e-01	medium
	Age	→	pc1a_Riesgo	0.194	0.184	1.270e-02	negligible
	Gender	→	pc1a_Riesgo	0.194	0.194	8.282e-05	negligible
	Education	→	pc1a_Riesgo	0.194	0.194	-3.548e-05	negligible

In conclusion, the results of the structural model analysis are graphically presented in Figure 5.1. The diagram reports the estimated path coefficients for the second-order specification. Asterisks indicate the level of statistical significance

obtained through bootstrapping procedures. The R^2 values for each endogenous construct are also displayed within the model, providing a visual summary of the explanatory power of the specification.

Figure 5.1: Structural model results



5.1.4. Mediation analysis

The mediation analysis examines whether the relationship between Smartphone Addiction (SA) and cybervictimization risk is transmitted through the L-RAT framework and Gambling behaviours, both independently and sequentially. The analysis is conducted using bootstrapping procedures, and statistical significance is evaluated through bias-corrected (BC) 95% confidence intervals. Results are reported in Table 5.10.

For each configuration, the total effect, direct effect, indirect effect and the Variance Accounted For (VAF) are estimated. The VAF indicates the proportion of the total effect explained by the indirect pathway. The classification of mediation follows Zhao et al. (2010), who distinguish between complementary mediation, competitive mediation, indirect-only mediation, direct-only non-mediation, and no-effect non-mediation.

When LRAT is introduced as mediator, two configurations are examined. In the first case (SA → LRAT → pc1a_Riesgo), both the direct and indirect effects are statistically significant. The VAF indicates that approximately 42% of the total effect of SA on cybervictimization risk is transmitted through LRAT. According to Zhao et al. (2010), this corresponds to complementary mediation, as both the direct and indirect

effects are significant and operate in the same direction. In the second case (Gambling → LRAT → pc1a_Riesgo), the indirect effect through LRAT is statistically significant, and the direct path from Gambling to cybervictimization risk also remains significant. The VAF is approximately 50%, indicating that about half of the total association operates through LRAT. This configuration also corresponds to complementary mediation, as defined by Zhao et al. (2010).

When Gambling is examined as mediator between SA and cybervictimization risk the indirect effect is statistically significant but comparatively small. The direct effect remains significant after including Gambling in the model. The VAF is slightly above 5%, indicating that only a limited portion of the total association is transmitted through this pathway. According to Zhao et al. (2010), this configuration corresponds to complementary mediation, as both direct and indirect effects coexist and are statistically significant, although the magnitude of the indirect component is modest.

A sequential mediation structure is tested in which SA influences Gambling, which in turn affects LRAT, ultimately predicting cybervictimization risk. In this case, the sequential indirect effect is statistically significant, although its magnitude is comparatively small. The direct effect remains statistically significant after including the chained mediation path. The VAF indicates that only a limited proportion of the total association is transmitted through this sequential mechanism. According to Zhao et al. (2010), this configuration also corresponds to complementary mediation, as both the indirect pathway and the direct path remain significant and operate in the same direction.

Table 5.10: Mediation analysis results

Total Sample (N = 7,064)						
Mediator	Total effect		Direct effect		Indirect effect	
LRAT	SA → pc1a_Riesgo		SA → pc1a_Riesgo		SA → LRAT → pc1a_Riesgo	
	Coefficient	BC CI 95%	Coefficient	BC CI 95%	Coefficient	BC CI 95% VAF
	0.260 ***	[0.235 ; 0.284]	0.151 ***	[0.125 ; 0.175]	0.109 ***	[0.097 ; 0.122] 42.062%
	Gambling → pc1a_Riesgo		Gambling → pc1a_Riesgo		Gambling → LRAT → pc1a_Riesgo	
	Coefficient	BC CI 95%	Coefficient	BC CI 95%	Coefficient	BC CI 95% VAF
	0.124 ***	[0.093 ; 0.155]	0.062 ***	[0.034 ; 0.090]	0.062 ***	[0.051 ; 0.074] 49.999%
Gambling	SA → pc1a_Riesgo		SA → pc1a_Riesgo		SA → Gambling → pc1a_Riesgo	
	Coefficient	BC CI 95%	Coefficient	BC CI 95%	Coefficient	BC CI 95% VAF
	0.159 ***	[0.134 ; 0.183]	0.151 ***	[0.125 ; 0.175]	0.008 **	[0.005 ; 0.013] 5.342%
Gambling & LRAT	SA → pc1a_Riesgo		SA → pc1a_Riesgo		SA → Gambling → LRAT → pc1a_Riesgo	
	Coefficient	BC CI 95%	Coefficient	BC CI 95%	Coefficient	BC CI 95% VAF
	0.159 ***	[0.134 ; 0.183]	0.151 ***	[0.125 ; 0.175]	0.009 ***	[0.006 ; 0.011] 5.342%

Note: The following convention values have been adopted for significance levels:

- t-values below 1.96 are considered not significant at the 5% level ($p > 0.05$)
- values between 1.96 and 2.58 are marked with * ($p < 0.05$)
- values between 2.58 and 3.29 are marked with ** ($p < 0.01$)
- values above 3.29 are marked with *** ($p < 0.001$)

5.2. Machine Learning results

To complement the structural equation modelling approach, a machine learning model is implemented to assess predictive performance and to identify the most relevant predictors of cybervictimization risk. The model is trained and evaluated following the procedure described in Chapter 4.

5.2.1. Model performance evaluation

The predictive performance of the model is assessed using a confusion matrix, standard classification metrics, and the Receiver Operating Characteristic (ROC) curve.

The classification results on the test set are reported in Table 5.11, which presents the confusion matrix computed using the default probability threshold of 0.5. The matrix shows the distribution of true positives, true negatives, false positives, and false negatives. The proportion of correctly classified observations reflects the overall predictive accuracy of the model, while the distribution of errors provides information on the balance between sensitivity and specificity.

Table 5.11: Confusion matrix

		Predicted	
		NonVictim	Victim
Actual	NonVictim	1041	28
	Victim	274	70

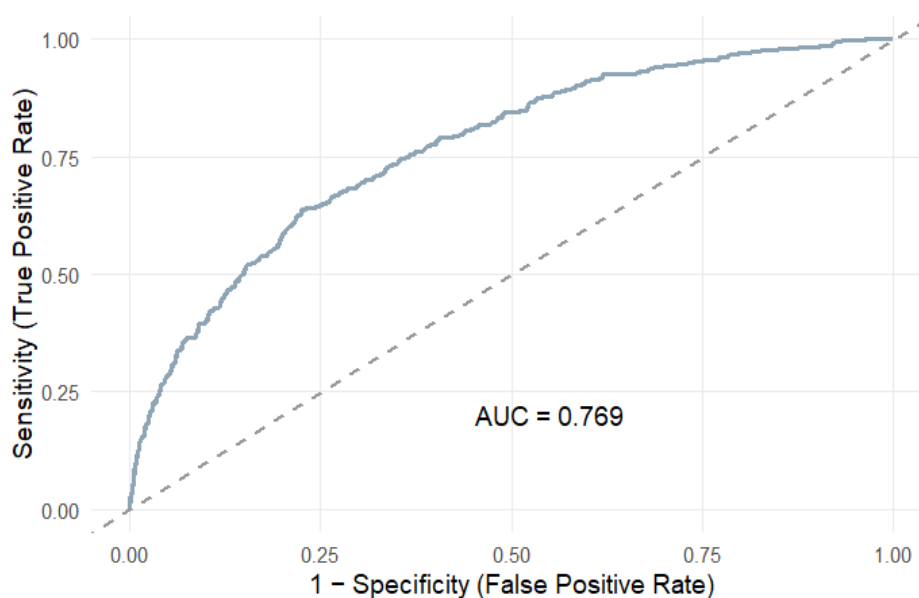
Table 5.12 reports the main classification metrics, including Accuracy, Sensitivity (Recall), Specificity, Precision , F1-score, and AUC. Although overall accuracy is relatively high (0.786), the model exhibits a strong imbalance between sensitivity and specificity at the default classification threshold of 0.5. In particular, the low sensitivity indicates that a substantial proportion of victimized individuals are not correctly identified, while specificity remains very high. This pattern suggests that the model is more effective in identifying non-victimized cases than in detecting high-risk individuals, a result that should be interpreted in light of the class distribution and threshold selection. The F1-score summarizes the trade-off between precision and recall, providing a balanced evaluation of classification performance in the presence of potential class imbalance. Finally, the AUC value of 0.769 confirms that the classifier performs better than random classification and demonstrates meaningful predictive capacity.

Overall, the reported metrics indicate satisfactory discriminative performance on the test set, although the imbalance between sensitivity and specificity highlights the relevance of threshold selection.

Table 5.12: Classification metrics

Metric	Estimator	Estimate
Accuracy	binary	0.786
Sensitivity (Recall)	binary	0.203
Specificity	binary	0.974
Precision (Positive Predictive Value)	binary	0.714
F1-score	binary	0.317
AUC	binary	0.769

The ROC curve displayed in Figure 5.2 illustrates the trade-off between the true positive rate and the false positive rate across different thresholds. Both the ROC curve and the associated AUC value confirm the model's discriminatory power.

Figure 5.2: ROC curve

5.2.2. Variable importance analysis

The relative contribution of predictors is evaluated through variable importance measures, reported in Table 5.13 and graphically represented in Figure 5.3. Variable importance is computed using permutation-based importance within the ranger implementation of Random Forests, quantifying the average decrease in predictive performance when each predictor is randomly permuted.

The analysis shows that Proximity emerges as the most influential variable, with a permutation importance of 2.207×10^{-2} . When permuted, this variable produces the largest reduction in classification accuracy, indicating a dominant role in the model's predictive mechanism. Smartphone Addiction follows with an importance of 1.552×10^{-2} , confirming its substantial predictive contribution within the Random Forest

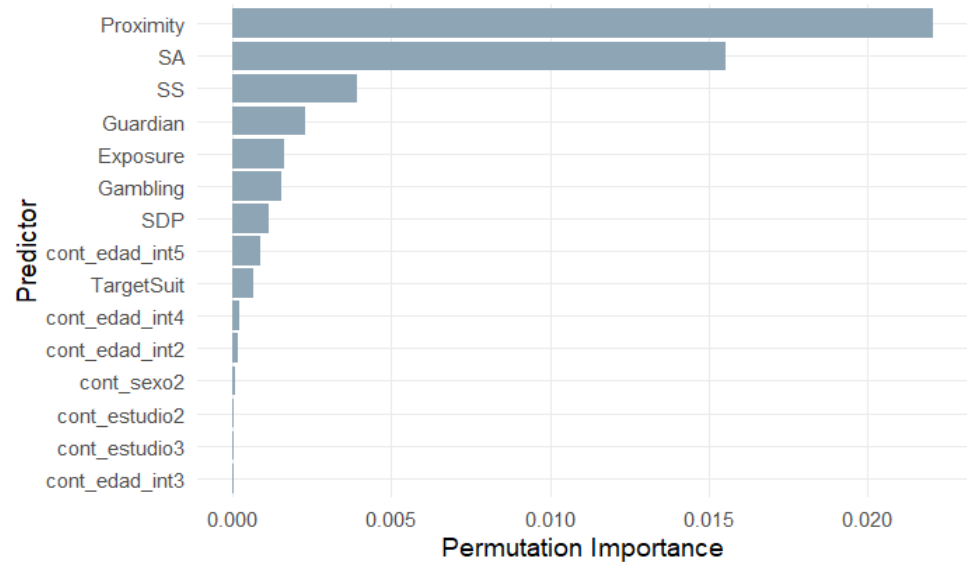
framework. Social Support ranks third, while Capable Guardianship, Exposure, and Gambling display more moderate importance values.

The remaining predictors, including SDP, Target Suitability, and the sociodemographic controls (age, gender, education), show comparatively lower importance scores, indicating limited incremental contribution to predictive accuracy relative to the top-ranked variables.

Table 5.13: Variable importance

Rank	Variable	Importance
1	Proximity	2.207e-02
2	SA	1.552e-02
3	SS	3.925e-03
4	Guardian	2.3027e-03
5	Exposure	1.644e-03
6	Gambling	1.567e-03
7	SDP	1.162e-03
8	cont_edad_int5	8.837e-04
9	TargetSuit	6.643e-04
10	cont_edad_int4	2.205e-04
11	cont_edad_int2	1.677e-04
12	cont_sexo2	8.040e-05
13	cont_estudio2	6.639e-05
14	cont_estudio3	6.159e-05
15	cont_edad_int3	3.371e-05

Figure 5.3: Variable importance plot



The distribution of variable importance values provides further insight into the structure of the predictive model. The minimum importance is approximately zero (slightly negative), indicating that some predictors contribute negligibly or introduce

minor noise. The median importance is approximately $7.740\text{e-}04$, suggesting that a substantial proportion of variables have only marginal predictive contribution. The mean importance ($\approx 3.148\text{e-}03$) exceeds the median, indicating a right-skewed distribution, where a limited number of predictors account for a disproportionately larger share of the model's predictive power. This pattern is confirmed by the maximum importance value ($\approx 2.207\text{e-}02$), which highlights the presence of a small subset of highly informative predictors relative to the remaining variables.

Overall, the distribution suggests that predictive performance is driven by a restricted group of variables, while most predictors contribute only minimally to classification accuracy.

5.3. Discussion

5.3.1. Convergence of explanatory and predictive evidence

The results of the PLS-SEM and Machine Learning analyses provide converging evidence regarding the determinants of cybervictimization risk. While the structural model evaluates theory-driven relationships and mediation mechanisms, the Random Forest model assesses predictive performance and variable relevance in a data-driven framework.

Across both approaches, variables derived from the Lifestyle-Routine Activity Theory (L-RAT) framework emerge as central. In the structural model, the higher-order L-RAT construct shows a statistically significant association with cybervictimization risk, supporting H1. In the machine learning model, Proximity to motivated offenders ranks as the most influential predictor in terms of permutation importance. When permuted, this variable produces the largest reduction in classification accuracy, indicating a dominant role in the predictive mechanism.

This convergence reinforces the criminological interpretation of cybervictimization as a function of routine activities and situational exposure rather than solely individual-level traits. Consistent with prior applications of L-RAT to cyberspace (Herrero, Torres, Vivas, Hidalgo, et al., 2021), the findings confirm that online behavioural patterns shape the convergence of motivated offenders, suitable targets, and guardianship conditions in digital environments.

5.3.2. L-RAT as central mechanism of cybervictimization

The structural results provide empirical support for the applicability of Lifestyle-Routine Activity Theory to online victimization processes. The significant direct association between L-RAT components and cybervictimization risk confirms H1, indicating that routine activity conditions are positively related to risk exposure in the digital context.

Moreover, mediation analysis shows that L-RAT partially mediates both the relationship between Smartphone Addiction and cybervictimization (H5) and the relationship between Gambling and cybervictimization (H3). In both cases, the indirect effects are statistically significant and classified as complementary mediation according to Zhao et al. (2010), as direct and indirect effects coexist and operate in the same direction. These findings align with prior research suggesting that behavioural patterns influence cyber-risk by altering daily routines and exposure settings (Herrero, Torres, Vivas, Hidalgo, et al., 2021). The mediation evidence strengthens this interpretation: addictive or gambling-related behaviours do not only exert direct effects but also reshape routine activities in ways that increase vulnerability.

Therefore, H1, H3, and H5 are supported, confirming that LRAT functions as a central explanatory mechanism in the model.

5.3.3. Smartphone addiction and behavioural vulnerability

Smartphone Addiction exhibits a statistically significant direct association with cybervictimization risk, supporting H4. This finding is consistent with prior literature linking problematic smartphone use to increased exposure to risky online environments and diminished self-regulatory control (as presented in Chapter 2).

The mediation results further indicate that a substantial portion of the association between Smartphone Addiction and cybervictimization risk operates indirectly through L-RAT (H5 supported). This suggests that addictive use patterns modify online routines, increasing proximity to potential offenders and reducing effective guardianship. By contrast, the mediation pathway through Gambling (H6) is statistically significant but modest in magnitude. The VAF indicates that only a small proportion of the total effect of Smartphone Addiction on cybervictimization operates through gambling behaviours. Although the indirect effect exists and is classified as complementary mediation, its contribution is limited relative to the LRAT-mediated pathway.

These results suggest that Smartphone Addiction represents a central behavioural vulnerability factor that operates primarily through altered routine activity patterns rather than exclusively through gambling behaviours.

5.3.4. Gambling behaviours and cyber risk

Online gambling behaviours show a statistically significant direct association with cybervictimization risk, supporting H2. This finding is consistent with prior evidence linking gambling activities to increased impulsivity, sensation-seeking tendencies, and exposure to high-risk digital environments.

Furthermore, the mediation analysis indicates that L-RAT partially mediates the gambling-cybervictimization relationship (H3 supported). Approximately half of the total association operates through routine activity mechanisms, indicating that

gambling contributes to cyber risk not only directly but also by modifying online exposure patterns.

This reinforces the interpretation of gambling as a behavioural context that amplifies routine activity risks, rather than as an isolated risk factor detached from situational mechanisms.

5.3.5. Psychosocial antecedents of Smartphone Addiction

The structural model partially supports the hypothesized psychosocial antecedents of Smartphone Addiction. Consistent with H8, perceived Social Support shows a strong and statistically significant negative association with Smartphone Addiction. This finding aligns with prior longitudinal evidence indicating that higher levels of social support function as protective factors against problematic smartphone use (Herrero, Torres, et al., 2019).

By contrast, the relationship between Social Digital Pressure and Smartphone Addiction does not support H7. Although the path coefficient is statistically significant, the effect is negative, indicating that higher levels of perceived digital social pressure are associated with slightly lower levels of Smartphone Addiction within this sample. This result diverges from previous findings suggesting a positive association between digital social pressure and problematic smartphone use (Herrero, Torres, Vivas, Arenas, et al., 2021).

The magnitude of the coefficient is small ($\beta = -0.058$), suggesting that Social Digital Pressure plays a limited role relative to Social Support. Nevertheless, the unexpected direction of the relationship calls for cautious interpretation. One possible explanation is that higher perceived digital pressure may activate compensatory regulatory mechanisms, reducing excessive engagement rather than amplifying it. Alternatively, contextual or sample-specific characteristics may moderate the association between perceived pressure and addictive behaviours.

5.3.6. The role of demographic factors

Control paths indicate that age is the only demographic covariate consistently associated with cybervictimization risk, while gender and education do not show statistically significant effects in the estimated specification.

From a modelling standpoint, this suggests that the main substantive relationships are not merely artefacts of demographic composition, at least with respect to the included covariates. At the same time, non-significant controls should not be interpreted as evidence of no demographic heterogeneity in general; rather, they indicate that within this model and sample, the variance attributable to these covariates is limited relative to the behavioural and routine-activity mechanisms.

5.3.7. Integrating criminological theory and behavioural addiction

Overall, the findings support a criminological interpretation of cybervictimization grounded in L-RAT while incorporating behavioural addiction as a mechanism that alters routine activity patterns. The results indicate that cybervictimization risk emerges from the interaction between situational exposure (L-RAT dimensions) and behavioural vulnerabilities (Smartphone Addiction and Gambling). The machine learning evidence strengthens this interpretation by independently identifying Proximity and Smartphone Addiction among the most influential predictors.

Thus, the study confirms the centrality of L-RAT in explaining cybervictimization in online gambling contexts, while also demonstrating that behavioural addiction processes act as upstream determinants that reshape routine activities and increase exposure to digital risk

6 Conclusions, limitations and future investigation

6.1. Conclusions

The present study examined how online gambling behaviour, smartphone addiction, and digital routines shape the likelihood of experiencing cybervictimization, grounding the analysis in Lifestyle-Routine Activity Theory and integrating behavioural and psychosocial factors within a unified empirical framework.

The results provide consistent evidence that cybervictimization risk cannot be interpreted as the outcome of isolated technical vulnerabilities. Rather, it emerges from the interaction between behavioural dispositions, digital practices, and situational exposure mechanisms embedded in everyday routines. By modelling Lifestyle-Routine Activity Theory (L-RAT) as a second-order formative construct composed of Exposure to motivated offenders, Proximity to motivated offenders, Target Suitability, and Capable Guardianship, the study confirms the centrality of opportunity structures in explaining cybercrime victimization in digital environments.

Across model specifications, Smartphone Addiction (SA) emerges as a significant and robust predictor of cybervictimization risk. The effect is both direct and indirect, operating through increased exposure to motivated offenders and reduced effective guardianship. This finding supports the theoretical assumption that addictive patterns of smartphone use reshape digital routines, intensifying online presence and amplifying contact with potentially harmful environments. Cyber risk, therefore, is not merely a function of time spent online, but of how behavioural dysregulation alters situational risk configurations.

Online gambling behaviour further contributes to this configuration. Gambling is positively associated with routine-based risk dimensions within the structural model, particularly those related to exposure patterns that increase the likelihood of contact with motivated offenders. This association suggests that engagement in gambling-related digital environments may intensify opportunity structures embedded in everyday online routines, resulting in a more vulnerable digital risk profile. Gambling platforms and related digital ecosystems can reasonably be interpreted as high-risk contexts in which opportunity convergence becomes more

frequent. In this sense, gambling appears to operate as an amplifier of routine-based exposure mechanisms rather than as an isolated determinant of victimization.

The integration of Social Digital Pressure (SDP) and Social Support (SS) into the model clarifies additional psychosocial pathways. Social Support displays a protective function, showing a negative association with cybervictimization risk and contributing to a more resilient digital behavioural profile. This finding reinforces the interpretation of social resources as informal guardianship mechanisms within the framework. By contrast, the relationship between Social Digital Pressure and Smartphone Addiction does not follow the initially hypothesized direction. The structural model suggests an inverse association, indicating that perceived digital pressure does not necessarily translate into addictive engagement patterns within this dataset. As discussed in Chapter 5, this result may reflect a measurement-related or contextual dynamic: individuals experiencing higher digital pressure may develop more conscious self-regulatory strategies, or the construct may capture normative social expectations rather than compulsive behavioural escalation.

From a methodological standpoint, the combined use of PLS-SEM and Machine Learning strengthens the robustness of these conclusions. The structural equation model provides explanatory clarity, identifying statistically significant relationships and mediation mechanisms. The Machine Learning models, by contrast, confirm the relative importance of the same key predictors in an out-of-sample predictive framework. The convergence between explanatory and predictive evidence reinforces confidence in the theoretical model and reduces concerns about overfitting or model-specific artefacts.

Overall, the study contributes to the literature in three main ways. First, it empirically validates the relevance of Lifestyle-Routine Activity Theory in digitally mediated contexts through a higher-order formative specification. Second, it demonstrates that behavioural addiction mechanisms, particularly Smartphone Addiction, constitute a structurally significant antecedent of cybervictimization risk. Third, it shows that explanatory modelling and predictive analytics can be successfully integrated to obtain a more comprehensive understanding of digital risk behaviours.

In increasingly digitised societies, cybervictimization must therefore be understood as a behavioural-situational phenomenon shaped by routine activities, psychosocial pressures, and platform-specific ecosystems. The findings highlight the need to move beyond purely technical prevention strategies and toward interventions that address behavioural regulation, digital literacy, and the restructuring of risky online routines.

6.2. Policy, practical, and social implications

The findings of this study carry relevant implications at policy, institutional, and societal levels. By demonstrating that cybervictimization risk emerges from the interaction between behavioural patterns, psychosocial conditions, and situational exposure mechanisms, the results challenge prevention strategies that focus exclusively on technical safeguards. Effective cybercrime prevention requires a broader behavioural and contextual approach.

From a public policy perspective, the empirical support for Lifestyle-Routine Activity Theory in digital environments suggests that interventions should target opportunity structures rather than solely individual awareness. Exposure to motivated offenders and weaknesses in capable guardianship represent key structural dimensions of risk. Policies aimed at strengthening digital guardianship (through default security configurations, mandatory authentication standards, and simplified privacy controls) may reduce vulnerability at scale without relying entirely on individual behavioural change.

The results also highlight the relevance of smartphone addiction as a behavioural vulnerability factor. Prevention programs should therefore integrate digital well-being initiatives within broader cybersecurity strategies. Educational campaigns addressing problematic smartphone use, compulsive engagement patterns, and excessive digital immersion may indirectly contribute to reducing exposure to cyber threats. In this sense, cybersecurity and digital health policies should not be treated as separate domains but as interconnected areas of intervention.

Online gambling environments represent an additional area of concern. While gambling is not presented as an autonomous determinant of victimization, its association with routine-based exposure patterns suggests that regulatory authorities should carefully monitor the digital ecosystems surrounding gambling platforms. Risk warnings, transparent communication practices, and strengthened platform-level safeguards may mitigate the amplification of exposure mechanisms identified in the model. Public health approaches to gambling harm could therefore incorporate cybersecurity awareness components.

The protective role of Social Support further underlines the importance of social resources in reducing digital vulnerability. Informal guardianship mechanisms, including peer advice, family support, and access to reliable information networks, appear to buffer cyber risk. Community-based initiatives and digital literacy programs should therefore encourage collective rather than purely individualistic approaches to online safety. Strengthening social capital may contribute indirectly to improved cybersecurity outcomes.

At an organisational level, the integration of explanatory and predictive modelling demonstrates the value of combining theory-driven analysis with data-

driven approaches. Institutions responsible for monitoring cyber risk, including public agencies and private platforms, may benefit from hybrid analytical frameworks capable of identifying both structural determinants and high-risk behavioural profiles. The convergence between PLS-SEM and Machine Learning findings suggests that theoretically grounded variables retain predictive relevance in applied risk detection systems.

Finally, at a societal level, the study contributes to reframing cybervictimization as a socially embedded phenomenon. Digital risks are shaped by everyday routines, platform architectures, and psychosocial pressures. Prevention efforts should therefore promote balanced digital practices, responsible platform design, and accessible security infrastructures. Reducing cybercrime requires structural alignment between individual behaviour, technological environments, and regulatory frameworks.

Overall, the implications of this research support the development of integrated prevention strategies that combine behavioural regulation, platform accountability, and institutional safeguards, moving beyond purely reactive or technically isolated responses to cyber threats.

6.3. Limitations and future research

Despite the robustness of the analytical strategy and the consistency of the results, several limitations should be acknowledged. First, the study relies on a cross-sectional research design. Although the structural model specifies directional relationships grounded in theory, the empirical evidence is based on contemporaneous observations. As a consequence, causal inferences cannot be formally established. The associations identified between Smartphone Addiction, routine-based exposure mechanisms, and cybervictimization risk should therefore be interpreted as structural relationships rather than definitive causal effects. Second, the analysis is based on self-reported survey data. Constructs such as Smartphone Addiction, online gambling behaviour, and cybervictimization experiences are measured through subjective responses, which may be affected by recall bias, social desirability bias, or misclassification of past incidents. Although the measurement model demonstrates satisfactory reliability and validity, the exclusive reliance on self-reported indicators may limit the precision of behavioural estimates. Third, the Machine Learning component is restricted to a single algorithmic approach. Although the Random Forest model provides stable predictive performance and convergent variable importance patterns, the absence of alternative algorithms limits the comparative assessment of predictive robustness. Finally, the dataset is derived from a national survey conducted within a specific regulatory and socio-cultural context. The generalizability of the findings to other countries or digital ecosystems with

different gambling regulations, cybersecurity infrastructures, or social norms cannot be automatically assumed.

These limitations open several directions for future research. Longitudinal designs would allow researchers to examine the temporal sequencing of behavioural addiction, changes in digital routines, and subsequent victimization risk. Panel data or repeated-measure designs could clarify dynamic processes and strengthen causal interpretation. Future studies should also conduct formal multi-group analyses and moderation tests to assess whether structural relationships differ across gamblers and non-gamblers, demographic subgroups, or varying levels of digital engagement. Such analyses would refine the understanding of heterogeneous risk profiles. From a predictive standpoint, subsequent research could compare alternative Machine Learning algorithms to evaluate the stability of classification performance and variable relevance across modelling strategies. Ensemble approaches or additional tree-based and regression-based models may provide complementary evidence regarding predictive accuracy. Finally, integrating survey data with behavioural indicators (e.g.: objective measures of device usage, app interaction logs, or transaction records) would enhance measurement precision and reduce reliance on subjective reporting. The combination of self-reported psychosocial constructs with behavioural data could produce a more comprehensive and ecologically valid representation of cybervictimization risk.

Overall, addressing these methodological and analytical extensions would contribute to consolidating the behavioural-situational framework proposed in this study and to advancing the empirical understanding of digital risk processes.

Time planning and budgeting

Time Planning

The development of this Master Thesis was structured into five main phases: Planning Phase, Evaluation of Alternative Theories, Model Estimation and Conduct of the Study, Writing Phase, and Closing Phase. This structure reflects the methodological progression from conceptual development to empirical modelling and final validation.

The Planning Phase focused on aligning the research with the objectives of the CIBERADICJUEGO2024 project and reviewing the available dataset. During this stage, the literature on Lifestyle-Routine Activity Theory, Smartphone Addiction, and online gambling was systematically analysed. In parallel, PLS-SEM and Machine Learning methodologies were examined in depth in order to define the integrated analytical strategy. The overall structure of the thesis was also outlined.

The Evaluation of Alternative Theories phase involved the assessment of alternative theoretical frameworks and the construction of potential complementary model specifications. This phase ensured that the final conceptual model was theoretically grounded and methodologically justified.

The Model Estimation and Conduct of the Study phase represented the core empirical stage. It included dataset inspection, variable mapping, specification of the PLS-SEM model, estimation of the structural relationships, and validation of the final model through robustness checks and bootstrapping procedures.

During Writing Phase, the theoretical framework, methodology, results, discussion, and conclusions were drafted. The explanatory findings from PLS-SEM and the predictive insights from Machine Learning were integrated into a coherent analytical narrative.

Finally, the Closing Phase was dedicated to comprehensive revision, formatting according to institutional guidelines, final validation of the manuscript, and submission.

Figure 7.1 shows the Gantt chart of the project.

Budgeting

The present Master Thesis was developed over a period of 34 weeks, from July 2025 to February 2026. Considering an estimated average dedication of 18 hours per week, the total research effort amounts to:

$$34 \text{ weeks} \times 18 \text{ hours/week} = 612 \text{ hours}$$

The estimated hourly compensation for the research activity is set at 12 €/hour, representing a conservative academic research rate.

Supervisory contribution is estimated following the reference structure adopted in comparable academic works. A total of 25 hours of supervision is considered, with an estimated hourly cost of 24 €/hour.

Regarding technical resources, the computer used for the development of the thesis has an estimated purchase price of 2,000 €, with an expected useful life of 8 years. Since the project duration is approximately 8 months (July 2025 – February 2026), the proportional amortization corresponds to:

$$(8 \text{ months} / 96 \text{ months}) \times 2,000 \text{ €} = 166.67 \text{ €}$$

Concerning software tools, both R and RStudio are open-source and free of charge. Microsoft 365 is provided through the university institutional license; therefore, no additional cost is attributed to software licenses.

The estimated cost structure is summarized in Table 7.1.

Table 7.1: *Estimated cost structure*

Category	Quantity	Unit Cost	Total Cost
Research work	612 h	12 €/h	7,344 €
Supervisory contribution (estimated)	25 h	24 €/h	600 €
Personal computer amortization	8 months (of 96)	2,000 €	166.67 €
Software tools	-	0 €	0 €
Total Estimated Budget			8,110.67 €

Ethical, legal, and professional considerations

The present study is conducted within the framework of the research projects “CIBERADICJUEGO2024” (Reference SUBV24/00012) and “The commodification of human behaviour: Data sharing and privacy threats among online gamblers” (Reference SUBV25/00014) and adheres to established ethical, legal, and professional standards applicable to academic research involving human-subject data and behavioural analysis.

From an ethical perspective, the study relies on survey data that were previously collected under the protocols of the broader research project. Participation in the original survey was voluntary, and respondents provided informed consent prior to data collection. The dataset used in this thesis does not contain personally identifiable information, and all analyses were conducted on anonymized data. No attempt was made to re-identify participants or to access sensitive personal information beyond the variables included in the dataset. The constructs analysed in this study, including Smartphone Addiction, online gambling behaviour, and cybervictimization experiences, involve potentially sensitive behavioural dimensions. For this reason, the analysis was conducted with particular attention to responsible interpretation. The results are presented in aggregate form and do not allow the identification of individuals or specific subgroups. The study avoids stigmatizing interpretations and does not frame vulnerable behavioural profiles as deterministic or pathological categories. From a legal standpoint, the research complies with applicable data protection regulations, including the General Data Protection Regulation (GDPR) of the European Union. Since the dataset was provided within an institutional research project, data handling procedures followed established institutional standards for secure storage, access limitation, and restricted processing. No external data sources were merged with the dataset, and no automated decision-making systems were implemented using individual-level predictions. Professional responsibility in this study is reflected in methodological transparency and analytical rigor. All modelling decisions, including the specification of the second-order formative construct, the estimation procedures in PLS-SEM, and the configuration of the Machine Learning model, are explicitly documented. Analytical steps are reproducible, and the statistical procedures are consistent with established methodological guidelines. The integration of explanatory and predictive approaches was implemented with the objective of strengthening robustness rather than maximizing statistical significance. Finally, the

interpretation of results has been conducted with caution regarding causal claims. Given the cross-sectional nature of the data, conclusions are limited to structural associations rather than definitive causal inferences. This restraint reflects adherence to responsible scientific practice.

Overall, the study respects ethical research principles, complies with relevant legal requirements, and aligns with professional standards of academic integrity, transparency, and methodological rigor.

Contribution to the sustainable development goals (SDGs)

The present study contributes to several United Nations Sustainable Development Goals (SDGs) by addressing the behavioural, technological, and institutional dimensions of cybervictimization risk in digitally mediated environments.

SDG 3 - Good Health and Well-being

The findings are directly connected to SDG 3, particularly in relation to mental health and behavioural well-being. The study identifies Smartphone Addiction as a significant predictor of increased exposure to cybervictimization risk. By modelling the interaction between addictive digital engagement and routine-based exposure mechanisms, the research highlights how problematic patterns of technology use may indirectly compromise psychosocial well-being.



Understanding the structural relationship between behavioural dysregulation and digital vulnerability supports preventive strategies aimed at reducing both cyber harm and associated psychological distress.

SDG 4 - Quality Education

The results also contribute to SDG 4 through their implications for digital literacy and education. The identification of capable guardianship as a protective factor emphasizes the importance of knowledge, awareness, and informed digital practices in reducing cyber risk. Strengthening digital education programs, particularly those addressing cybersecurity behaviours, responsible smartphone use, and online risk awareness, may mitigate the exposure mechanisms identified in the structural model.



The study therefore supports the integration of cybersecurity competence and digital self-regulation skills into educational initiatives.

SDG 9 - Industry, Innovation and Infrastructure

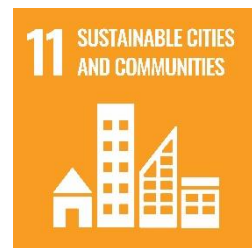
The research contributes to SDG 9 by providing empirical evidence relevant to the development of resilient digital infrastructures. By validating Lifestyle-Routine Activity Theory in a digitally mediated context, the study underscores the importance of secure platform design and institutional guardianship mechanisms.



The integration of explanatory (PLS-SEM) and predictive (Machine Learning) modelling also illustrates how advanced analytical methods can support innovation in risk detection and prevention strategies, strengthening the evidence base for digital infrastructure policies.

SDG 11 - Sustainable Cities and Communities

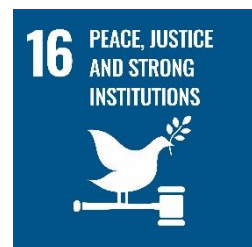
Digital environments increasingly constitute an integral component of contemporary communities. Cybervictimization affects not only individuals but also the broader digital ecosystem in which social, economic, and civic interactions occur. By identifying structural drivers of cyber risk, the study contributes to the creation of safer and more resilient digital communities.



Reducing vulnerability to online fraud and cybercrime enhances trust in digital interactions and supports the sustainability of digitally connected communities.

SDG 16 - Peace, Justice and Strong Institutions

Finally, the study contributes to SDG 16 by addressing cybercrime prevention and institutional resilience. Understanding how behavioural patterns, digital routines, and guardianship mechanisms interact to produce vulnerability provides insights relevant to public policy and institutional regulation.



Strengthening institutional responses to cybercrime requires not only technical solutions but also behavioural awareness and evidence-based policy interventions. By clarifying the determinants of cybervictimization risk, this research supports efforts aimed at reducing digital harm and reinforcing trust in online systems.

Overall, the study aligns with multiple SDGs by integrating behavioural science, cybersecurity research, and digital governance. It highlights the need for coordinated action across education, technological infrastructure, public policy, and individual behaviour to promote a secure and sustainable digital society.

References

- Akdemir, N., & Lawless, C. J. (2020). Exploring the human factor in cyber-enabled and cyber-dependent crime victimisation: A lifestyle routine activities approach. *Internet Research*, 30(6), 1665–1687. <https://doi.org/10.1108/INTR-10-2019-0400>
- American Psychiatric Association. (2022). *Diagnostic and Statistical Manual of Mental Disorders, Fifth Edition (DSM-5-TR)*. American Psychiatric Association Publishing. <https://psychiatryonline.org/doi/book/10.1176/appi.books.9780890425787>
- Anderson, C. L., & Agarwal, R. (2010). Practicing Safe Computing: A Multimethod Empirical Examination of Home Computer User Security Behavioral Intentions. *MIS Quarterly*, 34(3), 613–643. <https://doi.org/10.2307/25750694>
- Bagozzi, R. P., & Phillips, L. W. (1982). Representing and Testing Organizational Theories: A Holistic Construal. *Administrative Science Quarterly*, 27(3), 459–489. <https://doi.org/10.2307/2392322>
- Barroso, C., Carrión, G. C., & Roldán, J. L. (2010). Applying Maximum Likelihood and PLS on Different Sample Sizes: Studies on SERVQUAL Model and Employee Behavior Model. In V. Esposito Vinzi, W. W. Chin, J. Henseler, & H. Wang (Eds), *Handbook of Partial Least Squares* (pp. 427–447). Springer Berlin Heidelberg. https://doi.org/10.1007/978-3-540-32827-8_20
- Beyens, I., Frison, E., & Eggermont, S. (2016). “I don’t want to miss a thing”: Adolescents’ fear of missing out and its relationship to adolescents’ social needs, Facebook use, and Facebook related stress. *Computers in Human Behavior*, 64, 1–8. <https://doi.org/10.1016/j.chb.2016.05.083>
- Bian, M., & Leung, L. (2015). Linking loneliness, shyness, smartphone addiction symptoms, and patterns of smartphone use to social capital. *Social Science Computer Review*, 33(1), 61–79. <https://doi.org/10.1177/0894439314528779>
- Boehmke, B., & Greenwell, B. (2019). *Hands-On Machine Learning with R* (1st edn). Chapman and Hall/CRC. <https://doi.org/10.1201/9780367816377>
- Breiman, L. (2001). Random Forests. *Machine Learning*, 45(1), 5–32. <https://doi.org/10.1023/A:1010933404324>
- Busch, P. A., & McCarthy, S. (2021). Antecedents and consequences of problematic smartphone use: A systematic literature review of an emerging research area. *Computers in Human Behavior*, 114, 106414. <https://doi.org/10.1016/j.chb.2020.106414>
- Canale, N., Griffiths, M. D., Vieno, A., Siciliano, V., & Molinaro, S. (2016). Impact of Internet gambling on problem gambling among adolescents in Italy: Findings from a large-scale nationally representative survey. *Computers in Human Behavior*, 57, 99–106. <https://doi.org/10.1016/j.chb.2015.12.020>
- Choi, K.-S. (2008). Computer Crime Victimization and Integrated Theory: An Empirical Assessment. *International Journal of Cyber Criminology*, 2.
- Cohen, J. (1988). *Statistical Power Analysis for the Behavioral Sciences* (2nd edn). Routledge. <https://doi.org/10.4324/9780203771587>

- Cohen, L. E., & Felson, M. (1979). Social Change and Crime Rate Trends: A Routine Activity Approach. *American Sociological Review*, 44(4), 588–608. <https://doi.org/10.2307/2094589>
- Cohen, S., & Wills, T. A. (1985). Stress, social support, and the buffering hypothesis. *Psychological Bulletin*, 98(2), 310–357. <https://doi.org/10.1037/0033-2909.98.2.310>
- Cybersecurity in Spain: Challenges and Opportunities*. (n.d.). Springboard35. Retrieved 8 January 2026, from https://springboard35.com/en/blog/cybersecurity-in-spain-challenges-opportunities/?utm_source=chatgpt.com
- Dhir, A., Yossatorn, Y., Kaur, P., & Chen, S. (2018). Online social media fatigue and psychological wellbeing—A study of compulsive use, fear of missing out, fatigue, anxiety and depression. *International Journal of Information Management*, 40, 141–152. <https://doi.org/10.1016/j.ijinfomgt.2018.01.012>
- Diamantopoulos, A., Riefler, P., & Roth, K. P. (2008). Advancing formative measurement models. *Journal of Business Research, Formative Indicators*, 61(12), 1203–1218. <https://doi.org/10.1016/j.jbusres.2008.01.009>
- Diamantopoulos, A., & Winklhofer, H. M. (2001). Index Construction with Formative Indicators: An Alternative to Scale Development. *Journal of Marketing Research*, 38(2), 269–277. <https://doi.org/10.1509/jmkr.38.2.269.18845>
- Efron, B., & Tibshirani, R. J. (1994). *An Introduction to the Bootstrap*. Chapman and Hall/CRC. <https://doi.org/10.1201/9780429246593>
- Encord. (n.d.). *What is a Confusion Matrix?* | *Machine Learning Glossary* | Encord | Encord. Encord. Retrieved 29 January 2026, from <https://encord.com/glossary/confusion-matrix/>
- Erdem, H. D., Herrero, J., & Urueña, A. (2025). The Bidirectional Relationships between Social Pressure in Digital Contexts, Depression, and Social Support over Time. *Psychosocial Intervention*, 34(3), 189–200. <https://doi.org/10.5093/pi2025a15>
- Escario, J.-J., & Wilkinson, A. V. (2020). Exploring predictors of online gambling in a nationally representative sample of Spanish adolescents. *Computers in Human Behavior*, 102, 287–292. <https://doi.org/10.1016/j.chb.2019.09.002>
- Field, A. (2024). *Discovering Statistics Using IBM SPSS Statistics*. SAGE Publications.
- Fornell, C., & Larcker, D. F. (1981). Evaluating structural equation models with unobservable variables and measurement error. *Journal of Marketing Research*, 18(1), 39–50. <https://doi.org/10.2307/3151312>
- Gainsbury, S. M. (2015). Online Gambling Addiction: The Relationship Between Internet Gambling and Disordered Gambling. *Current Addiction Reports*, 2(2), 185–193. <https://doi.org/10.1007/s40429-015-0057-8>
- Gainsbury, S. M., Russell, A., Hing, N., Wood, R., Lubman, D., & Blaszczynski, A. (2015). How the Internet is Changing Gambling: Findings from an Australian Prevalence Survey. *Journal of Gambling Studies*, 31(1), 1–15. <https://doi.org/10.1007/s10899-013-9404-7>
- Garofalo, J. (1986). Lifestyles and Victimization: An Update. In E. A. Fattah (Ed.), *From Crime Policy to Victim Policy: Reorienting the Justice System* (pp. 135–155). Palgrave Macmillan UK. https://doi.org/10.1007/978-1-349-08305-3_7
- Gui, M., & Büchi, M. (2021). From Use to Overuse: Digital Inequality in the Age of Communication Abundance. *Social Science Computer Review*, 39(1), 3–19. <https://doi.org/10.1177/0894439319851163>

- Hadlington, L. (2017). Human factors in cybersecurity; examining the link between Internet addiction, impulsivity, attitudes towards cybersecurity, and risky cybersecurity behaviours. *Heliyon*, 3(7). <https://doi.org/10.1016/j.heliyon.2017.e00346>
- Haenlein, M., & Kaplan, A. M. (2004). A Beginner's Guide to Partial Least Squares Analysis. *Understanding Statistics*, 3(4), 283–297. https://doi.org/10.1207/s15328031us0304_4
- Hair, J. F., Hult, G. T. M., Ringle, C. M., Sarstedt, M., Danks, N. P., & Ray, S. (2021). *Partial Least Squares Structural Equation Modeling (PLS-SEM) Using R: A Workbook*. Springer International Publishing. <https://doi.org/10.1007/978-3-030-80519-7>
- Hair, J. F., Risher, J. J., Sarstedt, M., & Ringle, C. M. (2019). When to use and how to report the results of PLS-SEM. *European Business Review*, 31(1), 2–24. <https://doi.org/10.1108/EBR-11-2018-0203>
- Hair, J. F., Sarstedt, M., Ringle, C. M., & Mena, J. A. (2012). An assessment of the use of partial least squares structural equation modeling in marketing research. *Journal of the Academy of Marketing Science*, 40(3), 414–433. <https://doi.org/10.1007/s11747-011-0261-6>
- Hair, J., Ringle, C., & Sarstedt, M. (2011). PLS-SEM: Indeed a silver bullet. *The Journal of Marketing Theory and Practice*, 19, 139–151. <https://doi.org/10.2753/MTP1069-6679190202>
- Hastie, T., Tibshirani, R., & Friedman, J. (2009). *The Elements of Statistical Learning*. Springer. <https://doi.org/10.1007/978-0-387-84858-7>
- Haug, S., Castro, R. P., Kwon, M., Filler, A., Kowatsch, T., & Schaub, M. P. (2015). Smartphone use and smartphone addiction among young people in Switzerland. *Journal of Behavioral Addictions*, 4(4), 299–307. <https://doi.org/10.1556/2006.4.2015.037>
- Hayes, A. F. (2017). *Introduction to Mediation, Moderation, and Conditional Process Analysis, Second Edition: A Regression-Based Approach*. Guilford Publications.
- Henseler, J. (2018). Partial least squares path modeling: Quo vadis? *Quality & Quantity*, 52(1), 1–8. <https://doi.org/10.1007/s11135-018-0689-6>
- Henseler, J., Ringle, C. M., & Sarstedt, M. (2015). A new criterion for assessing discriminant validity in variance-based structural equation modeling. *Journal of the Academy of Marketing Science*, 43(1), 115–135. <https://doi.org/10.1007/s11747-014-0403-8>
- Herrero, J., Torres, A., Vivas, P., Arenas, Á. E., & Urueña, A. (2021). Examining the Empirical Links Between Digital Social Pressure, Personality, Psychological Distress, Social Support, Users' Residential Living Conditions, and Smartphone Addiction. *Social Science Computer Review*. (Sage CA: Los Angeles, CA). <https://doi.org/10.1177/0894439321998357>
- Herrero, J., Torres, A., Vivas, P., Hidalgo, A., Rodríguez, F. J., & Urueña, A. (2021). Smartphone Addiction and Cybercrime Victimization in the Context of Lifestyles Routine Activities and Self-Control Theories: The User's Dual Vulnerability Model of Cybercrime Victimization. *International Journal of Environmental Research and Public Health*, 18(7), 3763. <https://doi.org/10.3390/ijerph18073763>
- Herrero, J., Torres, A., Vivas, P., & Urueña, A. (2019). Smartphone Addiction and Social Support: A Three-year Longitudinal Study. *Psychosocial Intervention*, 28(3), 111–118. <https://doi.org/10.5093/pi2019a6>
- Herrero, J., Torres, A., Vivas, P., & Urueña, A. (2022). Smartphone Addiction, Social Support, and Cybercrime Victimization: A Discrete Survival and Growth Mixture Model. *Psychosocial Intervention*, 31(1), 59–66. <https://doi.org/10.5093/pi2022a3>

- Herrero, J., Urueña, A., Torres, A., & Hidalgo, A. (2019a). Smartphone addiction: Psychosocial correlates, risky attitudes, and smartphone harm. *Journal of Risk Research*, 22(1), 81–92. <https://doi.org/10.1080/13669877.2017.1351472>
- Herrero, J., Urueña, A., Torres, A., & Hidalgo, A. (2019b). Socially Connected but Still Isolated: Smartphone Addiction Decreases Social Support Over Time. *Social Science Computer Review*, 37(1), 73–88. <https://doi.org/10.1177/0894439317742611>
- Hindelang, M. J., Gottfredson, M. R., & Garofalo, J. (with Internet Archive). (1978). *Victims of personal crime: An empirical foundation for a theory of personal victimization*. Cambridge, Mass.: Ballinger Pub. Co. <http://archive.org/details/victimsofpersona0000hind>
- Holt & Bossler. (2015). *Cybercrime in Progress: Theory and prevention of technology-enabled offenses*. Routledge. <https://doi.org/10.4324/9781315775944>
- Hosmer, D. W., Lemeshow, S., & Sturdivant, R. X. (2013). *Applied Logistic Regression*. John Wiley & Sons.
- INRIA. (2025). *Machine Learning in Python with scikit-learn* [Online course materials, Zenodo, 2022]. <https://doi.org/10.5281/zenodo.7220306>.
- Instituto Nacional de Estadística. (2021). *Estadística continua de población. Datos definitivos a 1 de julio de 2021* [Data set]. <https://www.ine.es/>
- International Telecommunication Union. (2025). *Measuring digital development, Facts and Figures 2025*.
- Islam, M. A., & Bhuiyan, M. R. I. (2022). *Digital Transformation and Society* (SSRN Scholarly Paper No. 4604376). Social Science Research Network. <https://doi.org/10.2139/ssrn.4604376>
- James, G., Witten, D., Hastie, T., & Tibshirani, R. (2021). *An Introduction to Statistical Learning: With Applications in R*. Springer US. <https://doi.org/10.1007/978-1-0716-1418-1>
- James, R. J. E., Dixon, G., Dragomir, M.-G., Thirlwell, E., & Hitcham, L. (2023). Understanding the construction of ‘behavior’ in smartphone addiction: A scoping review. *Addictive Behaviors*, 137, 107503. <https://doi.org/10.1016/j.addbeh.2022.107503>
- Jarvis, C. B., MacKenzie, S. B., & Podsakoff, P. M. (2003). A Critical Review of Construct Indicators and Measurement Model Misspecification in Marketing and Consumer Research. *Journal of Consumer Research*, 30(2), 199–218. <https://doi.org/10.1086/376806>
- Joreskog, K. G. (1971). Simultaneous factor analysis in several populations. *Psychometrika*, 36(4), 409–426. <https://doi.org/10.1007/BF02291366>
- Kairouz, S., Paradis, C., & Nadeau, L. (2012). Are Online Gamblers More At Risk Than Offline Gamblers? *Cyberpsychology, Behavior, and Social Networking*, 15(3), 175–180. <https://doi.org/10.1089/cyber.2011.0260>
- Kramer, S. (2025). *Understanding online threat victimization: The roles of online activity, digital protective measures, gender and age* [Master Thesis]. <https://studenttheses.uu.nl/handle/20.500.12932/49629>
- Kutner, M. H. (2005). *Applied Linear Statistical Models*. McGraw-Hill Irwin.
- Leukfeldt, E. R., & Yar, M. (2016). Applying Routine Activity Theory to Cybercrime: A Theoretical and Empirical Analysis. *Deviant Behavior*, 37(3), 263–280. <https://doi.org/10.1080/01639625.2015.1012409>
- Lin, N., Ensel, W. M., & Vaughn, J. C. (1981). Social Resources and Strength of Ties: Structural Factors in Occupational Status Attainment. *American Sociological Review*, 46(4), 393. <https://doi.org/10.2307/2095260>

- Livingstone, C., & Rintoul, A. (2020). Moving on from responsible gambling: A new discourse is needed to prevent and minimise harm from gambling. *Public Health, Gambling: An Emerging Public Health Challenge*, 184, 107–112. <https://doi.org/10.1016/j.puhe.2020.03.018>
- Macaulay, P., Steer, O. L., & Betts, L. R. (2020). Factors leading to cyber victimization. In *Emerging Cyber Threats and Cognitive Vulnerabilities* (pp. 1–25). Academic Press. <https://www.sciencedirect.com/science/chapter/edited-volume/abs/pii/B9780128162033000010>
- Memon, M. A., Cheah, J.-H., Ramayah, T., Ting, H., & Chuah, F. (2018). MEDIATION ANALYSIS: ISSUES AND RECOMMENDATIONS. *Journal of Applied Structural Equation Modeling*, 2(1), i–ix. [https://doi.org/10.47263/JASEM.2\(1\)01](https://doi.org/10.47263/JASEM.2(1)01)
- Miethe, T. D., & Meier, R. F. (1994). *Crime and its Social Context: Toward an Integrated Theory of Offenders, Victims, and Situations*. SUNY Press.
- Miró-Llinares, F., Drew, J., & Townsley, M. (2020). *Understanding Target Suitability in Cyberspace: An International Comparison of Cyber Victimization Processes*. <https://doi.org/10.5281/ZENODO.3744874>
- Mitchell, T. M. (1997). *Machine Learning*. McGraw-Hill Education.
- Nahar, H., & Bhuyain, A. K. M. (2021). *Applying Routine Activity Theory to Cyber Victimization: A Secondary Analysis*.
- Organisation for Economic Co-operation and Development. (2019). *How's Life in the Digital Age?: Opportunities and Risks of the Digital Transformation for People's Well-being* How's Life in the Digital Age? OECD.
- Panova, T., & Carbonell, X. (2018). Is smartphone addiction really an addiction? *Journal of Behavioral Addictions*, 7(2), 252–. DOAJ Directory of Open Access Journals PubMed Central. <https://doi.org/10.1556/2006.7.2018.49>
- Piva, A. (2019, March 28). Cyber Crime: Cos'è, come affrontarlo e difendersi in azienda. *Osservatori Digital Innovation del Politecnico di Milano*. <https://www.osservatori.net/blog/cyber-security/cybercrime-definizione-italia/>
- Pratt, T. C., Holtfreter, K., & Reisig, M. D. (2010). Routine Online Activity and Internet Fraud Targeting: Extending the Generality of Routine Activity Theory. *Journal of Research in Crime and Delinquency*, 47(3), 267–296. <https://doi.org/10.1177/0022427810365903>
- Preacher, K. J., & Hayes, A. F. (2008). Asymptotic and resampling strategies for assessing and comparing indirect effects in multiple mediator models. *Behavior Research Methods*, 40(3), 879–891. <https://doi.org/10.3758/BRM.40.3.879>
- Reer, F., Tang, W. Y., & Quandt, T. (2019). Psychosocial well-being and social media engagement: The mediating roles of social comparison orientation and fear of missing out. *New Media & Society*, 21(7), 1486–1505. <https://doi.org/10.1177/1461444818823719>
- Rigdon, E., Sarstedt, M., & Ringle, C. (2017). On Comparing Results from CB-SEM and PLS-SEM: Five Perspectives and Five Recommendations. *Marketing ZFP*, 39(39), 4–16. <https://doi.org/10.15358/0344-1369-2017-3>
- Rock, P. (2002). *On becoming a victim* (C. Hoyle & R. Wilson, Eds; pp. 1–22). Hart Publishing. <http://www.hartpub.co.uk>
- Saqr, M., & López-Pernas, S. (Eds). (2024). *Learning Analytics Methods and Tutorials: A Practical Guide Using R*. Springer Nature Switzerland. <https://doi.org/10.1007/978-3-031-54464-4>

- Sarstedt, M., Hair, J. F., Cheah, J.-H., Becker, J.-M., & Ringle, C. M. (2019). How to Specify, Estimate, and Validate Higher-Order Constructs in PLS-SEM. *Australasian Marketing Journal*, 27(3), 197–211. <https://doi.org/10.1016/j.ausmj.2019.05.003>
- Shmueli, G., Sarstedt, M., Hair, J. F., Cheah, J.-H., Ting, H., Vaithilingam, S., & Ringle, C. M. (2019). Predictive model assessment in PLS-SEM: Guidelines for using PLSpredict. *European Journal of Marketing*, 53(11), 2322–2347. <https://doi.org/10.1108/EJM-02-2019-0189>
- Stefanovics, E. A., & Potenza, M. N. (2022). Update on Gambling Disorder. *Psychiatric Clinics of North America, Addiction Psychiatry: Challenges and Recent Advances*, 45(3), 483–502. <https://doi.org/10.1016/j.psc.2022.04.004>
- Stevic, A. (2024). Under Pressure? Longitudinal Relationships between Different Types of Social Media Use, Digital Pressure, and Life Satisfaction. *Social Media + Society*, 10(1), 20563051241239282. <https://doi.org/10.1177/20563051241239282>
- Streukens, S., & Leroi-Werelds, S. (2016). Bootstrapping and PLS-SEM: A step-by-step guide to get more out of your bootstrap results. *European Management Journal, European Management Research Using Partial Least Squares Structural Equation Modeling (PLS-SEM)*, 34(6), 618–632. <https://doi.org/10.1016/j.emj.2016.06.003>
- Tenenhaus, M., Vinzi, V. E., Chatelin, Y.-M., & Lauro, C. (2005). PLS path modeling. *Computational Statistics & Data Analysis, Partial Least Squares*, 48(1), 159–205. <https://doi.org/10.1016/j.csda.2004.03.005>
- Testa, G., Ruiz-Iniesta, A., García, O., Tarragón, E., Soriano, V., Benedetti, E., Cerrai, S., Molinaro, S., Brand, M., Potenza, M. N., & Mestre-Bach, G. (2025). Cross-jurisdictional factors linked to gambling frequency in adolescents from 28 European countries: A machine learning approach. *Psychiatry Research*, 351, 116602. <https://doi.org/10.1016/j.psychres.2025.116602>
- Thoits, P. A. (2011). Mechanisms Linking Social Ties and Support to Physical and Mental Health. *Journal of Health and Social Behavior*, 52(2), 145–161. <https://doi.org/10.1177/0022146510395592>
- Ting, C. H., & Chen, Y. Y. (2020). Smartphone addiction. In *Adolescent Addiction* (pp. 215–240). Academic Press. <https://doi.org/10.1016/B978-0-12-818626-8.00008-6>
- Urueña López, A., Mateo, F., Navío-Marco, J., Martínez-Martínez, J. M., Gómez-Sanchís, J., Vila-Francés, J., & José Serrano-López, A. (2019). Analysis of computer user behavior, security incidents and fraud using Self-Organizing Maps. *Computers & Security*, 83, 38–51. <https://doi.org/10.1016/j.cose.2019.01.009>
- Vuorre, M., & Przybylski, A. K. (2024). A multiverse analysis of the associations between internet use and well-being. *Technology, Mind, and Behavior*, 5(2), 65–75. <https://doi.org/10.1037/tmb0000127>
- Walklate, S. (2011). *Reframing criminal victimization: Finding a place for vulnerability and resilience*. <https://journals.sagepub.com/doi/abs/10.1177/1362480610383452>
- Wardle, H., Reith, G., Langham, E., & Rogers, R. D. (2019). Gambling and public health: We need policy action to prevent harm. *BMJ*, 365, 11807. <https://doi.org/10.1136/bmj.11807>
- Waweru, T. (2025, May 29). *Mobile Data Statistics 2025: Global Usage Trends and Consumption*. Tridens. <https://tridens technology.com/mobile-data-statistics/>
- Wetzels, M., Odekerken-Schröder, G., & van Oppen, C. (2009). Using PLS Path Modeling for Assessing Hierarchical Construct Models: Guidelines and Empirical Illustration. *MIS Quarterly*, 33(1), 177–195. <https://doi.org/10.2307/20650284>
- Whitty, M. T. (2019). Predicting susceptibility to cyber-fraud victimhood. *Journal of Financial Crime*, 26(1), 277–292. <https://doi.org/10.1108/JFC-10-2017-0095>

- Wood, M. (2005). Bootstrapped Confidence Intervals as an Approach to Statistical Inference. *Organizational Research Methods*, 8(4), 454–470. <https://doi.org/10.1177/1094428105280059>
- Young, K. S. (1998). Internet Addiction: The Emergence of a New Clinical Disorder. *CyberPsychology & Behavior*, 1(3), 237–244. <https://doi.org/10.1089/cpb.1998.1.237>
- Zhao, X., Lynch Jr., J. G., & Chen, Q. (2010). Reconsidering Baron and Kenny: Myths and truths about mediation analysis. *Journal of Consumer Research*, 37(2), 197–206. <https://doi.org/10.1086/651257>
- Zhou, Z., & Cheng, Q. (2025). Effects of online and offline social support on problematic smartphone use among Chinese adolescents: Mediating role of needs frustration. *Behaviour & Information Technology*, 1–13. <https://doi.org/10.1080/0144929X.2025.2585328>

List of Figures

<i>Figure 2.1: Structural model</i>	23
<i>Figure 3.1: Summary of differences between types of measurement models (Source: Jarvis et al., 2003)</i>	32
<i>Figure 3.2: General structure of a PLS-SEM model (Source: J. F. Hair et al., 2021)</i>	33
<i>Figure 3.3: Reflective measurement model assessment procedure (adapted from J. F. Hair et al., 2021)</i>	35
<i>Figure 3.4: Discriminant validity assessment using the HTMT (Source: J. F. Hair et al. 2021)</i>	37
<i>Figure 3.5: Formative measurement model assessment procedure (adapted from J. F. Hair et al., 2021)</i>	38
<i>Figure 3.6: Structural model assessment procedure (adapted from J. F. Hair et al., 2021)</i>	40
<i>Figure 3.7: Different types of higher-order constructs (adapted from J. F. Hair et al., 2021)</i>	44
<i>Figure 3.8: Conceptual representation of a mediation model (adapted from J. F. Hair et al., 2021)</i>	46
<i>Figure 3.9: Mediation analysis procedure (adapted from Hair et al., 2021, based on Zhao et al., 2010)</i>	47
<i>Figure 3.10: Main phases of the Machine Learning lifecycle</i>	55
<i>Figure 3.11: Conceptual representation of the Random Forest algorithm</i>	57
<i>Figure 3.12: Confusion matrix and model's performance evaluation metrics (adapted from Encord, n.d.)</i>	60
<i>Figure 3.13: ROC and AUC of a perfect model (a) and random guessing (b)</i>	61
<i>Figure 5.1: Structural model results</i>	95
<i>Figure 5.2: ROC curve</i>	98
<i>Figure 5.3: Variable importance plot</i>	99
<i>Figure 7.1: Gantt chart of the project</i>	112

List of Tables

Table 2.1: Model hypothesis	22
Table 3.1: Evaluation of reflective measurement models	49
Table 3.2: Evaluation of formative measurement models	49
Table 3.3: Evaluation of the structural model	50
Table 3.4: Evaluation of higher-order constructs (HOC)	50
Table 3.5: Mediation analysis	50
Table 3.6: Comparison between PLS-SEM and Machine Learning approaches	51
Table 3.7: Evaluation of the Machine Learning model	62
Table 4.1: Cybervictimization risk - Descriptive statistics	68
Table 4.2: L-RAT framework - Descriptive statistics of Exposure to motivated offender	69
Table 4.3: L-RAT framework - Descriptive statistics of Proximity to motivated offender	70
Table 4.4: L-RAT framework - Descriptive statistics of Target suitability	71
Table 4.5: L-RAT framework - Descriptive statistics of Capable guardianship	73
Table 4.6: Online gambling - Descriptive statistics	74
Table 4.7: Smartphone addiction - Descriptive statistics	76
Table 4.8: Social digital pressure - Descriptive statistics	77
Table 4.9: Social support - Descriptive statistics	78
Table 4.10: Sociodemographic variables - Descriptive statistics	80
Table 4.11: Machine Learning dataset - Descriptive statistics	83
Table 5.1: Reflective measurement model results	88
Table 5.2: Reflective measurement model results: discriminant validity	89
Table 5.3: Formative measurement model results	90
Table 5.4: Structural Model results, first-order model	92
Table 5.5: Structural Model results, second-order model	92
Table 5.6: Explanatory power (R^2) results, first-order model	93
Table 5.7: Explanatory power (R^2) results, second-order model	93
Table 5.8: Effect size (f^2) results, first-order model	94
Table 5.9: Effect size (f^2) results, second-order model	94
Table 5.10: Mediation analysis results	96

Table 5.11: <i>Confusion matrix</i>	97
Table 5.12: <i>Classification metrics</i>	98
Table 5.13: <i>Variable importance</i>	99
Table 7.1: <i>Estimated cost structure</i>	113

List of Equations

<i>Equation 1: Coefficient of determination (R^2)</i>	41
<i>Equation 2: Effect size (f^2)</i>	42
<i>Equation 3: Mediation analysis - Indirect effect</i>	46
<i>Equation 4: Mediation analysis - Total effect</i>	46
<i>Equation 5: True Positive Rate (a) and False Positive Rate (b)</i>	60

List of acronyms and symbols

List of acronyms	
AUC	Area Under the Curve
AVE	Average Variance Extracted
BC	Bias-Corrected bootstrap interval
CI	Confidence Interval
CR	Composite Reliability
ETSII	Escuela Técnica Superior de Ingenieros Industriales
FN	False Negative
FP	False Positive
FPR	False Positive Rate
GDPR	General Data Protection Regulation
HOC	Higher-Order Construct
HTMT	Heterotrait–Monotrait ratio of correlations
L-RAT	Lifestyle-Routine Activity Theory
LOC	Lower-Order Construct
ML	Machine Learning
PLS-SEM	Partial Least Squares Structural Equation Modelling
POLIMI	Politecnico di Milano
POMP	Percent of Maximum Possible
PPV	Positive Predictive Value
RF	Random Forest
ROC	Receiver Operating Characteristic
SA	Smartphone Addiction
SDP	Social Digital Pressure
SDG	Sustainable Development Goal
SEM	Structural Equation Modelling
SS	Social Support
TFM	Trabajo de Fin de Máster (Master Thesis)

TN	True Negative
TNR	True Negative Rate
TP	True Positive
TPR	True Positive Rate
UPM	Universidad Politécnica de Madrid
VAF	Variance Accounted For
VIF	Variance Inflation Factor

List of symbols

α	Cronbach's alpha (Internal consistency reliability)
β	Standardized path coefficient (or weight for formative indicators)
ε	Structural error term
f^2	Effect size (Cohen's f-squared)
w	Formative indicator weight
λ	Indicator loading
η	Latent variable
p	p-value (Statistical significance level)
R^2	Coefficient of determination
ρ_A	Exact reliability (Dijkstra-Henseler's criterion)
ρ_C	Composite reliability (Jöreskog's criterion)
t	t-statistic (from bootstrapping)

A Appendix - R code

```
"TFM Trevisan Federico
Date: 2026-02-03
Title: Understanding cybervictimization through smartphone addiction, online
gambling and routine activities: an integrated approach using PLS-SEM and ML"

# =====
# 1. CLEANING, DIRECTORY AND LIBRARIES

rm(list = ls())

current_path <- dirname(rstudioapi::getActiveDocumentContext())$path)
setwd(current_path)

library(haven)
library(tidyverse)
library(dplyr)
library(psych)
library(purrr)
library(tibble)
library(seminr)
library(boot)
library(tidymodels)
library(yardstick)
library(car)
library(readxl)
library(caret)
library(vip)
library(ranger)

set.seed(123)

# =====
# 2. DATA IMPORT

# Importing "df_agg", starting database
load("dati_grande.RData")

# Convert "haven_labelled" in "numeric"
df_agg <- df_agg %>%
  mutate(across(where(~ inherits(.x, "haven_labelled")), as.numeric))

# RECODING SCALES
# Target Suitability (g2_pcan_1:5) -> inverted: ↑ = ↑Suitability
df_agg <- df_agg %>%
  mutate(across(starts_with("g2_pcan_"), ~ 6 - .x))

# Social Digital Pressure (i14pcan_1:3) -> inverted: ↑ = ↑SDP
df_agg <- df_agg %>%
```

```

mutate(across(starts_with("i14pcan_"), ~ 6 - .x))

# Social Support (i9pcan_1:3) -> inverted: ↑ = ↑SS
df_agg <- df_agg %>%
  mutate(across(starts_with("i9pcan_"), ~ 6 - .x))

# =====
# 3. DATABASE - PLS-SEM

# First-order Formative Constructs (FOC L-RAT)
exposure_vars <- paste0("a1_pcan_", c(6,7,8,13))
proximity_vars <- paste0("d1_pcan_", 1:7)
target_vars <- paste0("g2_pcan_", 1:5)
guardian_vars <- paste0("g3_pcan_", c(1,3,5,7,9))

# Smartphone Addiction (reflective)
add_vars <- c("i6aan_4", "i6aan_5", "i6aan_9", "i6aan_10",
             "i6ban_2", "i6ban_6", "i6ban_8", "i6ban_9")

# Social Digital Pressure (reflective)
sdp_vars <- c("i14pcan_1", "i14pcan_2", "i14pcan_3")

# 3.4 Social Support (reflective)
ss_vars <- c("i9pcan_1", "i9pcan_2", "i9pcan_3")

# Gambling
gambling <- "jugador"

# Dependent variable pc1a_Riesgo (0-6 continuous)
df_agg$pc1a_Riesgo <- rowSums(df_agg[, paste0("pc1a_pcan_", c(1,2,4,5,6,7))],
                             na.rm = TRUE)
output <- "pc1a_Riesgo"

# DATABASE - PLS-SEM (ex: TFM_mod)
DB_PLSSEM <- df_agg %>%
  select(all_of(c(exposure_vars, proximity_vars, target_vars, guardian_vars,
                 add_vars, sdp_vars, ss_vars, gambling, output,
                 "cont_sexo", "cont_edad_int", "cont_estudio")) %>%
  drop_na())

# =====
# 4. DATABASE - ML

# POMP SUM for FORMATIVE CONSTRUCTS
pomp_sum <- function(df, items, min_item, max_item){
  k <- length(items)
  raw_sum <- rowSums(df[,items])
  min_total <- k*min_item
  max_total <- k*max_item
  ((raw_sum-min_total)/(max_total-min_total)) * 100}

# POMP MEAN for REFLECTIVE CONSTRUCTS
pomp_mean <- function(df, items, min_item, max_item){
  m <- rowMeans(df[, items])
  ((m - min_item)/(max_item - min_item))*100}

DB_ML_PROV <- DB_PLSSEM %>%
  mutate(

```

```

# Formative constructs
Exposure = pomp_sum(., exposure_vars, min_item = 0, max_item = 1),
Proximity = pomp_sum(., proximity_vars, min_item = 0, max_item = 1),
TargetSuit = pomp_sum(., target_vars, min_item = 1, max_item = 5),
Guardian = pomp_sum(., guardian_vars, min_item = 1, max_item = 5),

# Reflective constructs
SA = pomp_mean(., add_vars, min_item = 1, max_item = 5),
SDP = pomp_mean(., sdp_vars, min_item = 1, max_item = 5),
SS = pomp_mean(., ss_vars, min_item = 1, max_item = 5),

# Target rescaled to 0-100
pc1a_Riesgo = round((pc1a_Riesgo / 6) * 100, 2),

# Gambling rescaled to 0-100 (0=non-gambler, 100=gambler)
Gambling = as.numeric(jugador) * 100
) %>%
transmute(
  Exposure, Proximity, TargetSuit, Guardian, SA, SDP, SS, Gambling,
  cont_sexo, cont_edad_int, cont_estudio, pc1a_Riesgo
)

# Ensure categorical controls are factors (for one-hot encoding)
DB_ML_PROV <- DB_ML_PROV %>%
  mutate(
    cont_sexo = factor(cont_sexo),
    cont_estudio = factor(cont_estudio),
    cont_edad_int = factor(cont_edad_int)
  )

# One-hot encoding ONLY for sociodemographic controls
dummies <- model.matrix(
  ~ cont_sexo + cont_estudio + cont_edad_int,
  data = DB_ML_PROV
)[, -1]
dummies <- as_tibble(dummies)

# Final ML dataset: scaled numeric predictors + target + dummies
DB_ML <- bind_cols(
  DB_ML_PROV %>%
    select(Exposure, Proximity, TargetSuit, Guardian, SA, SDP, SS, Gambling,
    pc1a_Riesgo) %>%
    mutate(across(where(is.numeric), ~ round(.x, 2))),
  dummies
)

# =====
# 5. PLS-SEM: FIRST-ORDER MODEL

# 5.1 - MEASUREMENT MODEL (first-order model)
mod_mm_base <- constructs(
  # FORMATIVES
  composite("Exposure", exposure_vars, weights = mode_B),
  composite("Proximity", proximity_vars, weights = mode_B),
  composite("TargetSuit", target_vars, weights = mode_B),
  composite("Guardian", guardian_vars, weights = mode_B),
  # REFLECTIVES
  composite("SA", add_vars, weights = mode_A),

```

```

composite("SDP", sdp_vars, weights = mode_A),
composite("SS", ss_vars, weights = mode_A),
# SINGLE-ITEM
composite("Gambling", single_item("jugador")),
composite("Age", single_item("cont_edad_int")),
composite("Gender", single_item("cont_sexo")),
composite("Education", single_item("cont_estudio")),
composite("Output", single_item("pc1a_Riesgo"))
)

# 5.2 - STRUCTURAL MODEL (first-order model)
mod_sm_base <- relationships(
  paths(from = c("SS", "SDP"), to = "SA"),
  paths(from = "SA", to = "Gambling"),
  paths(from = c("SA", "Gambling"),
    to = c("Exposure", "Proximity", "TargetSuit", "Guardian")),
  paths(from = c("SA", "Gambling", "Exposure", "Proximity", "TargetSuit", "Guardian",
    "Age", "Gender", "Education"),
    to = "Output")
)

# 5.3 - ESTIMATION OF PLS-SEM (first-order model)
mod_base <- estimate_pls(data = DB_PLSSEM,
  measurement_model = mod_mm_base,
  structural_model = mod_sm_base)

# 5.4 - Bootstrap of the model to obtain significance levels (first-order model)
boot_base <- bootstrap_model(mod_base, nboot = 10000)
summary(boot_base)

# 5.5 - Reliability measures for reflexive constructs (first-order model)
reliability_base <- summary(mod_base)$reliability
reliability_base

# 5.6 - Fornell-Larcker criterion ( $\sqrt{AVE}$  > inter-construct correlations)
cat("\nbase - sqrt(AVE):\n")
print(sqrt(reliability_base[c("SA", "SDP", "SS"), "AVE"]))

# Correlations between reflective constructs from latent scores
FL_lv_scores <- as.data.frame(mod_base$construct_scores)
cat("\nCorrelations between reflective constructs (SA, SDP, SS):\n")
print(round(cor(FL_lv_scores[, c("SA", "SDP", "SS")]), 3))

# 5.7 - Bias-corrected bootstrap Confidence interval (first-order model)
# (following indications of Streukens and Leroi-Werelds, 2016)
bc_excel_ci <- function(t_star, t0, conf = 0.95,
  index_round = c("round", "ceiling_floor")){
  # --- Step 1: select t0, t_star ---
  index_round <- match.arg(index_round)
  t_star <- as.numeric(t_star)
  t_star <- t_star[is.finite(t_star)]
  R <- length(t_star)

  # --- Step 2: acceleration coefficient a ( $MD^3 / MD^2$ ) ---
  m <- mean(t_star)
  md3_sum <- sum((t_star - m)^3)
  md2_sum <- sum((t_star - m)^2)
  denom <- 6 * (md2_sum^(3/2))

```

```

a      <- if (denom == 0) 0 else (md3_sum / denom)
# intermediate denominator D5 = 6 * ( (sum MD^2)^(3/2) )

# --- Step 3.1: z0 from bootstrap proportion ---
p_less <- mean(t_star < t0)
eps    <- 0.5 / R
p_less <- min(max(p_less, eps), 1 - eps)
z0     <- qnorm(p_less)
alpha  <- (1 - conf) / 2
zL     <- qnorm(alpha)
zU     <- qnorm(1 - alpha)
# Z-lower e Z-upper (95% -> +/-1.96; 90% -> +/-1.645)

# --- Step 3.2: z' (prime) ---
zL_prime <- z0 + (z0 + zL) / (1 - a * (z0 + zL))
zU_prime <- z0 + (z0 + zU) / (1 - a * (z0 + zU))

# --- Step 3.3: p-values associated to z' ---
pL_prime <- pnorm(zL_prime)
pU_prime <- pnorm(zU_prime)

# --- Step 4.1: observation numbers = p' * R ---
kL_raw <- pL_prime * R
kU_raw <- pU_prime * R
if (index_round == "round") {
  kL <- as.integer(round(kL_raw))
  kU <- as.integer(round(kU_raw))
}else{
  kL <- as.integer(ceiling(kL_raw))
  kU <- as.integer(floor(kU_raw))
}
# conservative choice: lower = ceiling, upper = floor
kL <- max(1L, min(R, kL))
kU <- max(1L, min(R, kU)) # clamp among 1 and R

# --- Step 4.2: bounds through statistical order (SMALL) ---
s      <- sort(t_star)
lower  <- s[kL]
upper  <- s[kU]
list(lower = lower, upper = upper, t0 = t0, mean_boot = m, R = R, a = a,
      z0 = z0, p_less = p_less, zL = zL, zU = zU, zL_prime = zL_prime,
      zU_prime = zU_prime, pL_prime = pL_prime, pU_prime = pU_prime,
      kL_raw = kL_raw, kU_raw = kU_raw, kL = kL, kU = kU)
}

bc_excel_ci_all_paths <- function(boot_obj, conf = 0.95, only_nonzero = TRUE,
                                  index_round = c("round", "ceiling_floor")){
  B0    <- boot_obj$path_coef
  BP    <- boot_obj$boot_paths
  preds <- rownames(B0)
  targs <- colnames(B0) # automatic path selection
  idx   <- which(!is.na(B0), arr.ind = TRUE)
  if (only_nonzero) idx <- idx[B0[idx] != 0, , drop = FALSE]
  R     <- dim(BP)[3]
  out   <- vector("list", nrow(idx))
  for (k in seq_len(nrow(idx))) {
    i    <- idx[k, 1]
    j    <- idx[k, 2]

```

```

    name    <- paste0(preds[i], " -> ", targs[j])
    t0      <- B0[i, j]
    t_star  <- BP[i, j, ]
    ci      <- bc_excel_ci(t_star = t_star, t0 = t0, conf = conf,
                          index_round = match.arg(index_round))
    out[[k]] <- data.frame(path = name, beta = ci$t0, bc_l = ci$lower,
                          bc_u = ci$upper, a = ci$a, z0 = ci$z0, kL = ci$kL,
                          kU = ci$kU, R = ci$R, row.names = NULL)
  }
  do.call(rbind, out)
}

tab_base <- bc_excel_ci_all_paths(boot_base, conf = 0.95, only_nonzero = TRUE,
                                index_round = "ceiling_floor")
tab_base

# 5.8 - R2 bias-corrected bootstrap CI (first-order model)
bootstrap_R2_from_paths_and_lvcor <- function(boot_obj, target, use_abs = TRUE){
  B0      <- boot_obj$path_coef # Beta original matrix (predictor x target)
  if (!(target %in% colnames(B0))) stop("target not a column in path_coef.")
  betas0 <- B0[, target]
  preds   <- names(betas0)[!is.na(betas0) & betas0 != 0]
  # predictors as incoming path != 0
  if (length(preds) == 0) stop("No incoming predictor for this target variable")
  C0      <- cor(boot_obj$construct_scores, use = "pairwise.complete.obs")
  # LV correlations from original model
  r_xy    <- C0[preds, target] # correlations pred->target vector
  if (use_abs) r_xy <- abs(r_xy)
  BP      <- boot_obj$boot_paths
  R       <- dim(BP)[3]
  # bootstrap path coefficients (array pred x target x R)
  # for each pred i, extract beta bootstrap towards the target: matrix R x nPred
  beta_star <- sapply(preds, function(p) BP[p, target, ])
  beta_star <- as.matrix(beta_star)
  if (use_abs) beta_star <- abs(beta_star)
  R2_star  <- as.numeric(beta_star %*% r_xy) # Eq.6: line product sum (replica)
  # original estimate, coherent with the formula
  beta0_vec <- betas0[preds]
  if (use_abs) beta0_vec <- abs(beta0_vec);
  R2_0     <- sum(beta0_vec * r_xy, na.rm = TRUE)
  list(target= target, preds= preds, lv_cor= r_xy, R2_0= R2_0, R2_star= R2_star)
}

all_endogenous <- function(B) colnames(B)[colSums(abs(B), na.rm = TRUE) > 0]

endo <- all_endogenous(boot_base$path_coef) # for each endogenous var

r2_results <- lapply(endo, function(y){
  tmp <- bootstrap_R2_from_paths_and_lvcor(boot_base, target = y, use_abs = TRUE)
  ci  <- bc_excel_ci(tmp$R2_star, tmp$R2_0, conf = 0.95)
  data.frame(target = y, R2_0 = tmp$R2_0, bc_l = ci$lower, bc_u = ci$upper,
             a = ci$a, z0 = ci$z0, R = ci$R, row.names = NULL)
})

r2_table <- do.call(rbind, r2_results)
r2_table

# 5.9 - Effect size (first-order model)

```

```

compute_f2_from_r2_df <- function(boot_obj, r2_df, endo_col = "target",
                                r2_col = "R2_0"){
  B0 <- boot_obj$path_coef
  C0 <- cor(boot_obj$construct_scores, use = "pairwise.complete.obs")
  out <- list()
  k <- 0
  for (i in seq_len(nrow(r2_df))){
    Y <- as.character(r2_df[[endo_col]][i])
    R2_incl <- as.numeric(r2_df[[r2_col]][i])
    if (!is.finite(R2_incl)) next
    if (!(Y %in% colnames(B0))) next
    betas <- B0[, Y]
    preds <- names(betas)[!is.na(betas) & betas != 0]
    for (X in preds) {beta_j <- betas[X]
      cor_j <- C0[X, Y]
      contrib_j <- beta_j * cor_j
      R2_excl <- R2_incl - contrib_j
      denom <- (1 - R2_incl)
      f2 <- if (denom <= 0) NA_real_ else (R2_incl - R2_excl) / denom
      k <- k + 1; out[[k]] <- data.frame(from = X, to = Y,
                                         R2_incl = R2_incl, R2_excl = R2_excl,
                                         f2 = f2, row.names = NULL)
    }
  }
  res <- do.call(rbind, out)
  res$size <- cut(res$f2, breaks = c(-Inf, 0.02, 0.15, 0.35, Inf),
                 # size effect label (Cohen)
                 labels = c("negligible", "small", "medium", "large"))
  res
}

f2_results <- compute_f2_from_r2_df(boot_obj = boot_base, r2_df = r2_table,
                                endo_col = "target", r2_col = "R2_0")
f2_results

# 5.10 - VIF first-order formative items
calc_vif <- function(data, indicators){
  purrr::map(indicators, function(v){
    others <- setdiff(indicators, v)
    if(length(others) == 0) return(NA)
    m <- lm(as.formula(paste(v, "~", paste(others, collapse = " + "))),
            data = data)
    tibble(Indicator = v, VIF = max(car::vif(m)))
  }) %>% bind_rows()

form_vif_base <- bind_rows(
  calc_vif(DB_PLSSSEM, exposure_vars) %>% mutate(Block = "Exposure"),
  calc_vif(DB_PLSSSEM, proximity_vars) %>% mutate(Block = "Proximity"),
  calc_vif(DB_PLSSSEM, target_vars) %>% mutate(Block = "TargetSuit"),
  calc_vif(DB_PLSSSEM, guardian_vars) %>% mutate(Block = "Guardian")
)

print(form_vif_base, n = Inf)

# 5.11 - INNER VIF (LOC model): collinearity among predictor constructs
calc_inner_vif <- function(model){
  # latent variable scores
  scores <- as.data.frame(model$construct_scores)

```

```

# path coefficient matrix: rows = predictors, cols = outcomes
B <- model$path_coef
# endogenous constructs: those with at least one incoming path
endo <- colnames(B)[colSums(abs(B), na.rm = TRUE) > 0]
# for each endogenous construct, compute VIF for its predictor set
res <- purrr::map_dfr(endo, function(y){
  preds <- rownames(B)[abs(B[, y, drop = TRUE]) > 0]
  preds <- setdiff(preds, y) # safety, should not happen
  # VIF only makes sense if at least 2 predictors
  if(length(preds) < 2){
    return(tibble(
      Outcome = y,
      Predictor = preds,
      VIF = NA_real_
    ))
  }
  m <- lm(as.formula(paste(y, "~", paste(preds, collapse = " + "))),
    data = scores)
  tibble(
    Outcome = y,
    Predictor = names(car::vif(m)),
    VIF = as.numeric(car::vif(m))
  )
})
res
}

inner_vif_base <- calc_inner_vif(mod_base)
print(inner_vif_base, n = Inf)

# =====
# 6. PLS-SEM: SECOND-ORDER MODEL

# 6.0 - Construction of LRAT as HOC
lv_scores_s1 <- as.data.frame(mod_base$construct_scores)

# Dataset final model
DB_PLSSEM_SECOND <- DB_PLSSEM %>%
  mutate(LV_Exposure = lv_scores_s1$Exposure,
    LV_Proximity = lv_scores_s1$Proximity,
    LV_TargetSuit = lv_scores_s1$TargetSuit,
    LV_Guardian = lv_scores_s1$Guardian )

# 6.1 - MEASUREMENT MODEL (second-order model)
mod_mm_second <- constructs(
  # FORMATIVE HOC
  composite("LRAT",
    multi_items("LV_", c("Exposure", "Proximity", "TargetSuit", "Guardian")),
    weights = mode_B),
  # REFLECTIVES
  composite("SA", add_vars, weights = mode_A),
  composite("SDP", sdp_vars, weights = mode_A),
  composite("SS", ss_vars, weights = mode_A),
  # SINGLE ITEMS
  composite("Gambling", single_item("jugador")),
  composite("Age", single_item("cont_edad_int")),
  composite("Gender", single_item("cont_sexo")),

```

```

    composite("Education", single_item("cont_estudio")),
    composite("Output",    single_item("pc1a_Riesgo"))
  )

# 6.2 - STRUCTURAL MODEL (second-order model)
mod_sm_second <- relationships(
  paths(from = c("SS","SDP"), to = "SA"),
  paths(from = "SA", to = "Gambling"),
  paths(from = c("SA","Gambling"), to = "LRAT"),
  paths(from = c("SA","Gambling","LRAT","Age","Gender","Education"),
        to = "Output")
)

# 6.3 - ESTIMATION OF PLS-SEM (second-order model)
mod_second <- estimate_pls(data = DB_PLSSEM_SECOND,
                          measurement_model = mod_mm_second,
                          structural_model  = mod_sm_second)

# 6.4 - Bootstrap of the model to obtain significance levels
boot_second <- bootstrap_model(mod_second,
                              nboot = 10000,
                              cores = parallel::detectCores())

summary(boot_second)

# 6.5 - OUTER VIF for HOC formative indicators (LV_ -> LRAT)
form_vif_hoc <- calc_vif(DB_PLSSEM_SECOND,
                        c("LV_Exposure", "LV_Proximity", "LV_TargetSuit", "LV_Guardian")) %>%
  mutate(Block = "LRAT (HOC)")

print(form_vif_hoc, n = Inf)

# 6.6 - INNER VIF (HOC model): collinearity among predictor constructs
inner_vif_second <- calc_inner_vif(mod_second)
print(inner_vif_second, n = Inf)

# 6.7 - Bias-corrected bootstrap Conf interval
tab_base_2 <- bc_excel_ci_all_paths(boot_second,
                                    conf = 0.95,
                                    only_nonzero = TRUE,
                                    index_round = "ceiling_floor")

tab_base_2

# 6.8 - R2 bias-corrected bootstrap CI (second-order model)
endo_bis <- all_endogenous(boot_second$path_coef)
r2_results_bis <- lapply(endo_bis,
                        function(y){
                          tmp <- bootstrap_R2_from_paths_and_lvcor(boot_second,
                                                                    target = y,
                                                                    use_abs = TRUE)
                          ci <- bc_excel_ci(tmp$R2_star, tmp$R2_0, conf = 0.95)
                          data.frame(target = y, R2_0 = tmp$R2_0,
                                      bc_l= ci$lower, bc_u = ci$upper, a = ci$a,
                                      z0 = ci$z0, R = ci$R, row.names = NULL)
                        })
r2_table_bis <- do.call(rbind, r2_results_bis)
r2_table_bis

# 6.9 - Effect size (second-order model)

```

```

f2_results_bis <- compute_f2_from_r2_df(boot_obj = boot_second,
                                       r2_df    = r2_table_bis,
                                       endo_col = "target", r2_col = "R2_0")

f2_results_bis

# 6.10 - Mediation Analysis (second-order model)
boot_t_value <- function(t0, t_star) {
  t0 / sd(t_star, na.rm = TRUE)
}

total_bc_ci <- function(boot_obj, x = "SA", m = "LRAT",
                        y = "Output", conf = 0.95) {
  # point estimates
  b_dir0 <- boot_obj$path_coef[x, y]
  b1_0   <- boot_obj$path_coef[x, m]
  b2_0   <- boot_obj$path_coef[m, y]
  ie_0   <- b1_0 * b2_0
  te_0   <- b_dir0 + ie_0
  # bootstrap draws
  b_dir_star <- as.numeric(boot_obj$boot_paths[x, y, ])
  b1_star    <- as.numeric(boot_obj$boot_paths[x, m, ])
  b2_star    <- as.numeric(boot_obj$boot_paths[m, y, ])
  ie_star    <- b1_star * b2_star
  te_star    <- b_dir_star + ie_star
  # BC CIs
  direct_ci  <- bc_excel_ci(b_dir_star, b_dir0, conf = conf)
  indirect_ci <- bc_excel_ci(ie_star,    ie_0,    conf = conf)
  total_ci   <- bc_excel_ci(te_star,    te_0,    conf = conf)
  # t-values
  dir_t <- boot_t_value(b_dir0, b_dir_star)
  ind_t <- boot_t_value(ie_0,    ie_star)
  tot_t <- boot_t_value(te_0,    te_star)
  # Output
  path_label <- paste(x, m, y, sep = " -> ")
  data.frame(
    Type = c(path_label, "Direct", "Indirect", "Total"),
    Estimate = c(NA, round(c(b_dir0, ie_0, te_0), 3)),
    Lower = c(NA, round(c(direct_ci$lower, indirect_ci$lower,
                          total_ci$lower), 3)),
    Upper = c(NA, round(c(direct_ci$upper, indirect_ci$upper,
                          total_ci$upper), 3)),
    t_stat = c(NA, round(c(dir_t, ind_t, tot_t), 3)),
    VAF = c(round(ie_0/te_0*100, 3), NA, NA, NA)
  )
}

# SA -> LRAT -> Output
med_SA_LRAT_Output <- total_bc_ci(boot_obj = boot_second, x = "SA",
                                  m = "LRAT", y = "Output", conf = 0.95)

med_SA_LRAT_Output
# Gambling -> LRAT -> Output
med_Gambling_LRAT_Output <- total_bc_ci(boot_obj = boot_second, x = "Gambling",
                                         m = "LRAT", y = "Output", conf = 0.95)

med_Gambling_LRAT_Output
# SA -> Gambling -> Output
med_SA_Gambling_Output <- total_bc_ci(boot_obj = boot_second, x = "SA",
                                       m = "Gambling", y = "Output", conf = 0.95)

med_SA_Gambling_Output

```

```

# Serial mediation function
serial_med_complete <- function(boot_obj, x, m1, m2, y, conf = 0.95) {
  # Original coefficients
  b_x_m1 <- boot_obj$path_coef[x, m1]
  b_m1_m2 <- boot_obj$path_coef[m1, m2]
  b_m2_y <- boot_obj$path_coef[m2, y]
  b_direct_0 <- boot_obj$path_coef[x, y]
  # Indirect and Total
  ind_ser_0 <- b_x_m1 * b_m1_m2 * b_m2_y
  tot_ser_0 <- ind_ser_0 + b_direct_0
  # Bootstrap extraction
  star1 <- as.numeric(boot_obj$boot_paths[x, m1, ])
  star2 <- as.numeric(boot_obj$boot_paths[m1, m2, ])
  star3 <- as.numeric(boot_obj$boot_paths[m2, y, ])
  star_direct <- as.numeric(boot_obj$boot_paths[x, y, ])
  ind_ser_star <- star1 * star2 * star3
  tot_ser_star <- ind_ser_star + star_direct
  # CI Bias-Corrected
  ci_ind <- bc_excel_ci(ind_ser_star, ind_ser_0, conf = conf)
  ci_dir <- bc_excel_ci(star_direct, b_direct_0, conf = conf)
  ci_tot <- bc_excel_ci(tot_ser_star, tot_ser_0, conf = conf)
  # Output
  path_label <- paste(x, m1, m2, y, sep = " -> ")
  data.frame(
    Type = c(path_label, "Direct", "Indirect", "Total"),
    Estimate = c(NA, round(c(b_direct_0, ind_ser_0, tot_ser_0), 3)),
    Lower = c(NA, round(c(ci_dir$lower, ci_ind$lower, ci_tot$lower), 3)),
    Upper = c(NA, round(c(ci_dir$upper, ci_ind$upper, ci_tot$upper), 3)),
    t_stat = c(NA, round(c(b_direct_0 / sd(star_direct),
                                ind_ser_0 / sd(ind_ser_star),
                                tot_ser_0 / sd(tot_ser_star)), 3)),
    VAF = c(round(ind_ser_0/tot_ser_0*100, 3), NA, NA, NA)
  )}
# SA -> Gambling -> LRAT -> Output
med_SA_G_LRAT_Output <- serial_med_complete(boot_second, "SA", "Gambling",
                                              "LRAT", "Output")
med_SA_G_LRAT_Output

# =====
# 7. ML - MODEL INITIALIZATION

# 7.0 - define binary target (classification)
DB_ML <- DB_ML %>%
  mutate(
    RiskBin = factor(ifelse(pcl1a_Riesgo > 0, "Victim", "NonVictim"),
                      levels = c("NonVictim", "Victim"))
  )
set.seed(123)

# 7.1 - Train/test split (80/20), stratified on the binary outcome
split_rf <- initial_split(DB_ML, prop = 0.80, strata = RiskBin)
train_data <- training(split_rf)
test_data <- testing(split_rf)

# 7.2 10-fold CV on training set
folds_rf <- vfold_cv(train_data, v = 10, strata = RiskBin)

# 7.3 - Recipe

```

```

rec_rf <- recipe(RiskBin ~ ., data = train_data) %>%
  # Removing pc1a_Riesgo
  step_rm(pc1a_Riesgo) %>%
  # Removing zero variance variables
  step_zv(all_predictors())

# 7.4 - Random Forest specification (classification)
rf_spec <- rand_forest(
  mode = "classification",
  trees = tune(), # to be optimised
  mtry = tune(), # to be optimised
  min_n = tune() # to be optimised
) %>%
  set_engine(
    "ranger",
    importance = "permutation", # measurements of varying importance
    sample.fraction = 0.632, # sampling fraction per tree (bootstrap ~ 63.2%)
    max.depth = tune(), # to be optimised
    replace = TRUE # sampling with replacement
  )

wf_rf <- workflow() %>%
  add_model(rf_spec) %>%
  add_recipe(rec_rf)

# 7.5 Parameter ranges
rf_params <- parameters(
  trees(range = c(500, 2000)),
  mtry() %>% finalize(select(train_data, -RiskBin, -pc1a_Riesgo)),
  min_n(range = c(2, 30)),
  tree_depth(range = c(2, 20))
)
rf_params$id[rf_params$id == "tree_depth"] <- "max.depth"

# 7.6 - Tuning grid
set.seed(123)
rf_grid <- grid_space_filling(rf_params, size = 30)
colnames(rf_grid)

# =====
# 8. ML - MODEL TUNING VIA CROSS-VALIDATION

set.seed(123)

# 8.1 - Model tuning
rf_tuned <- tune_grid(
  wf_rf,
  resamples = folds_rf,
  grid = rf_grid,
  metrics = metric_set(roc_auc, accuracy),
  control = control_grid(save_pred = TRUE)
)

# 8.2 - Selection
best_rf <- select_best(rf_tuned, metric = "roc_auc")
best_rf

wf_rf_final <- finalize_workflow(wf_rf, best_rf)

```

```

rf_final_fit <- last_fit(
  wf_rf_final,
  split = split_rf,
  metrics = metric_set(roc_auc, accuracy)
)

collect_metrics(rf_final_fit)

# =====
# 9. ML - EVALUATION ON THE TEST SET

# 9.1 - Collect test predictions from last_fit()
rf_test_preds <- collect_predictions(rf_final_fit)

# Sanity check: what columns are available?
dplyr::glimpse(rf_test_preds)

# 9.2 - Confusion matrix (default threshold 0.5)
cat("\n\nConfusion matrix (default threshold 0.5)\n")
conf_mat(rf_test_preds, truth = RiskBin, estimate = .pred_class)

# 9.3 - Full set of classification metrics
metrics_full <- metric_set(
  yardstick::accuracy,
  yardstick::sens,      # recall (TPR)
  yardstick::spec,      # specificity
  yardstick::ppv,        # precision
  yardstick::f_meas,
  yardstick::roc_auc
)

metrics_full(
  rf_test_preds,
  truth = RiskBin,
  estimate = .pred_class,
  .pred_Victim,
  event_level = "second"
)

# 9.4 - ROC curve data (for plotting/reporting)
roc_df <- roc_curve(
  rf_test_preds,
  truth = RiskBin,
  .pred_Victim,
  event_level = "second"
)

roc_df

# 9.5 - Plot ROC curve
autoplot(roc_df)

library(ggplot2)

auc_val <- yardstick::roc_auc(
  rf_test_preds,
  truth = RiskBin,

```

```

    .pred_Victim,
    event_level = "second"
  )$.estimate

ggplot(roc_df, aes(x = 1 - specificity, y = sensitivity)) +
  geom_line(linewidth = 1.2, color = "#8EA5B6") +
  geom_abline(
    linetype = "dashed",
    color = "grey60",
    linewidth = 0.8
  ) +
  annotate(
    "text",
    x = 0.55,
    y = 0.20,
    label = paste0("AUC = ", round(auc_val, 3)),
    size = 5
  ) +
  labs(
    x = "1 - Specificity (False Positive Rate)",
    y = "Sensitivity (True Positive Rate)"
  ) +
  theme_minimal(base_size = 14) +
  theme(
    panel.grid.minor = element_blank()
  )
)

# =====
# 10. ML - VARIABLE IMPORTANCE

# 10.1 - Final fit of the workflow
rf_final_workflow_full <- wf_rf_final %>%
  fit(data = DB_ML)

# 10.2 - Extraction of the ranger model
rf_engine <- extract_fit_engine(rf_final_workflow_full)

# 10.3 - Variable importance (permutation)
rf_var_imp <- ranger::importance(rf_engine)

rf_var_imp_tbl <- tibble::tibble(
  variable = names(rf_var_imp),
  importance = as.numeric(rf_var_imp)
) %>%
  dplyr::arrange(desc(importance))

rf_var_imp_tbl

summary(rf_var_imp_tbl$importance) # Sanity check

# 10.4 - Selection of "relevant" variables
rf_var_imp_top <- rf_var_imp_tbl %>%
  dplyr::filter(importance > 0) %>%
  dplyr::slice_max(importance, n = 15)

rf_var_imp_top

# 10.5 - Plot

```

```
library(ggplot2)

ggplot(rf_var_imp_top,
       aes(x = reorder(variable, importance), y = importance)) +
  geom_col(fill = "#8EA5B6") +
  coord_flip() +
  labs(
    x = "Predictor",
    y = "Permutation Importance"
  ) +
  theme_minimal(base_size = 14) +
  theme(
    panel.grid.minor = element_blank()
  )
```

