

Cyber Human Risk Framework: Measuring and Mitigating Human-Centered Cyber Risk

Troy C. Troublefield^{1,2,3} 

¹Department of Cyberpsychology, Capitol Technology University, Laurel, MD, USA

²Department of Information Technology, Capella University, Minneapolis, MN, USA

³Department of International Business, International School of Management, Paris, France

Email: drtroutroublefield@yahoo.com

How to cite this paper: Troublefield, T.C.

(2026) Cyber Human Risk Framework: Measuring and Mitigating Human-Centered Cyber Risk. *Journal of Information Security*, 17, 431-463.

<https://doi.org/10.4236/jis.2026.174020>

Received: May 7, 2026

Accepted: August 22, 2026

Published: August 25, 2026

Copyright © 2026 by author(s) and

Scientific Research Publishing Inc.

This work is licensed under the Creative

Commons Attribution International

License (CC BY 4.0).

<http://creativecommons.org/licenses/by/4.0/>



Open Access

Abstract

The proliferation of sophisticated cyber threats has increasingly exposed the inadequacy of purely technical defenses in protecting organizational and national security infrastructure. Human behavior remains the predominant attack surface exploited by adversaries across all sectors, yet existing cybersecurity frameworks insufficiently account for the psychological, cognitive, and behavioral dimensions of human-centered risk. This article introduces the Cyber Human Risk Framework (CHRF), an original theoretical model grounded in cyberpsychology, cognitive science, behavioral economics, and organizational psychology. The CHRF posits that human cyber risk is a dynamic, multi-layered construct composed of five interdependent domains: Cognitive Vulnerability, Behavioral Susceptibility, Affective Influence, Social Engineering Exposure, and Organizational Context. The framework introduces a novel risk quantification schema, the Human Cyber Risk Index (HCRI), enabling practitioners, researchers, and policymakers to assess, benchmark, and mitigate the human dimensions of cyber risk at individual, team, and enterprise levels. The CHRF draws upon established theoretical predecessors including the Technology Acceptance Model, Dual-Process Theory, Terror Management Theory, Social Influence Theory, and Cyberpsychology's unique contributions to understanding online behavior and identity. Theoretical synthesis, the empirical foundations drawn from existing behavioral and cyberpsychology research, and provisional practical application mechanisms are elaborated across a structured discussion that includes framework architecture, domain exposition, measurement considerations, and integration with existing cybersecurity governance standards. The CHRF's quantitative elements, including the Human Cyber Risk Index formula, domain weights, and risk tier thresholds, are explicitly provisional and are offered as empirical hypotheses requiring validation rather than as established measurement

parameters.

Keywords

Cyberpsychology, Cyber Human Risk, Behavioral Cybersecurity, Human Factors, Cognitive Vulnerability, Social Engineering, Organizational Security Culture, CHRF

1. Introduction

The contemporary cyber threat landscape is defined not by the sophistication of technical exploits alone but by the systematic exploitation of human psychology. Ransomware operators, nation-state adversaries, insider threats, and transnational criminal organizations have converged on a common strategic insight: the human element represents the most reliable and most exploitable vulnerability in any security architecture. According to a 2024 Verizon Data Breach Investigations Report, over 74% of all data breaches involved a human element, whether through error, privilege misuse, social engineering, or stolen credentials. This figure has remained disturbingly consistent over successive reporting cycles, underscoring the structural nature of the problem [1].

Despite this empirical clarity, the dominant discourse within cybersecurity continues to privilege technical solutions, endpoint detection, network monitoring, cryptographic protocols, and vulnerability patching, over the behavioral and psychological interventions that address the root of the problem. Frameworks such as the NIST Cybersecurity, and the Center for Internet Security Critical Security Controls (CIS Controls) offer organizational governance structures but treat human behavior largely as a compliance variable rather than a complex, dynamic, and scientifically tractable risk domain [2]-[4].

Cyberpsychology, the scientific study of the psychological phenomena associated with and influenced by emerging technology, offers a transformative lens through which to reconceptualize human cyber risk. Scholars including [5]-[7] have established that digital environments alter cognition, identity, social behavior, and decision-making in ways that both create vulnerabilities and offer intervention pathways. Yet the discipline has not produced a comprehensive, practitioner-deployable risk framework that translates cyberpsychology theory into structured, measurable, and governance-ready tools.

The Cyber Human Risk Framework (CHRF) is introduced in this article to address that gap. The CHRF is an original theoretical framework that synthesizes contributions from cognitive psychology, behavioral economics, organizational behavior, social influence theory, and cyberpsychology into a coherent, multi-domain model of human cyber risk. It advances beyond prior human-factor models by introducing structured domain architecture, a quantifiable risk index, and explicit linkages to governance standards, workforce development, and organiza-

tional culture change. The CHRF is positioned as a theoretically derived scholarly contribution to cyberpsychology and as a provisional practical instrument for applied cybersecurity risk management, one that offers structured assessment and governance architecture in advance of the empirical validation studies that will ultimately establish its measurement properties and predictive validity. Practitioners adopting the CHRF prior to validation should treat its quantitative outputs as directional indicators and governance-support tools rather than precision risk measurements.

This article proceeds through a structured exposition of the CHRF's theoretical foundations, domain architecture, measurement schema, and practical applications. Following the literature review, the five CHRF domains are elaborated in detail. The Human Cyber Risk Index is then introduced, followed by discussion of framework integration with existing governance standards, research implications, and limitations. The article concludes with recommendations for future empirical validation and applied implementation.

2. Methodology

2.1. Framework Development Approach and Methodological Status

The framework is grounded in structured narrative synthesis across five contributing disciplines. Rather than originating from a single primary data collection effort, the framework is grounded in systematic theoretical synthesis across five contributing disciplines, cognitive psychology, behavioral economics, organizational behavior, social influence theory, and cyberpsychology, assembled through an extensive review of peer-reviewed literature, empirical studies, and validated psychometric instruments spanning multiple decades of research. This approach reflects the foundational phase of framework development in which theoretical architecture must precede empirical validation, a sequence well-established in organizational and health behavioral science where models such as the Technology Acceptance Model, Protection Motivation Theory, and the Health Belief Model were similarly advanced as theoretically derived constructs prior to their subsequent empirical testing and refinement.

The literature review informing the CHRF's domain architecture was conducted with attention to empirical effect sizes, theoretical mechanism specificity, and direct relevance to security-relevant behavioral outcomes. Source selection prioritized peer-reviewed empirical studies documenting measurable relationships between psychological constructs and security behavior, foundational theoretical works providing the cognitive and behavioral mechanisms through which human cyber risk operates, and validated assessment instruments demonstrating psychometric adequacy in organizational and behavioral cybersecurity contexts. Particular weight was assigned to research demonstrating cross-study consistency of findings, as the CHRF's domain architecture is intended to reflect stable, replicable mechanisms rather than context-specific experimental artifacts.

Domain construction followed a deductive theoretical reasoning process in which established behavioral science constructs were mapped onto documented cybersecurity risk pathways. Each of the five domains, Cognitive Vulnerability, Behavioral Susceptibility, Affective Influence, Social Engineering Exposure, and Organizational Context, was derived from the convergence of multiple theoretical traditions addressing a distinct class of human security risk mechanisms. Domain boundaries were drawn to reflect theoretical distinctiveness and empirical separability, with particular attention to ensuring that each domain implicates different intervention targets and different assessment methodologies, thereby preserving the framework's practical utility for differentiated risk management rather than producing redundant constructs that collapse into a single underlying dimension.

The Human Cyber Risk Index was developed through a rational weighting methodology in which provisional domain weights were derived from the relative effect sizes and predictive validities documented in the behavioral cybersecurity literature. Cognitive Vulnerability was assigned a weight of 0.20, Behavioral Susceptibility 0.25, Affective Influence 0.15, Social Engineering Exposure 0.25, and Organizational Context 0.15, with all weights summing to 1.0 across a normalized 0 - 100 scale per domain. These weights are explicitly characterized as empirical hypotheses rather than validated parameters, representing the best available estimate from existing literature in the absence of original longitudinal outcome data directly linking domain scores to incident frequencies. The four-tier risk stratification schema, Low (0 - 24), Moderate (25 - 49), Elevated (50 - 74), and Critical (75 - 100), was derived from established risk stratification conventions in public health and organizational risk management, adapted to the cybersecurity context through alignment with industry-standard risk tier definitions employed in NIST SP 800-30 and FAIR risk quantification methodologies.

The organizational HCRI aggregation methodology, incorporating role-risk multipliers to produce the Role-Adjusted Organizational HCRI, was developed through analysis of documented adversarial targeting patterns drawn from threat intelligence reporting, insider threat research, and social engineering attack taxonomy literature. Role-risk multipliers were assigned based on three dimensions: access privilege level, organizational decision-making authority, and empirically documented adversarial targeting frequency for specific role categories. This methodology extends established role-based access control principles from technical security architecture into the human risk assessment domain, applying the logic of differential access risk to the assessment of differential social engineering exposure and organizational impact potential.

Framework integration with existing governance standards, NIST CSF 2.0, ISO/IEC 27001:2022, and CISA insider threat program guidelines, was accomplished through systematic crosswalk analysis in which each CHRF domain and assessment output was mapped against the specific control categories, risk assessment requirements, and documentation obligations of each standard. This cross-

walk methodology follows the alignment conventions established by the NIST National Cybersecurity Center of Excellence and is consistent with the approach taken in prior human-factor cybersecurity framework development efforts including the HAIS-Q instrument and the Security Culture Framework. The intent was not to demonstrate compliance but to identify the specific governance contexts in which CHRF outputs generate direct, documentable value, thereby reducing adoption friction for organizations operating within existing compliance architectures.

The CHRF's validation requirements, articulated as a forward research agenda rather than completed methodology, follow the sequential validation framework established in applied measurement science. The recommended validation sequence proceeds from content validity, expert review of domain completeness and theoretical representativeness, through construct validity testing using structural equation modeling to assess the proposed five-factor architecture against alternative competing models, to criterion validity testing in which HCRI scores are examined for predictive relationships with documented security incident outcomes in organizational settings. Test-retest reliability assessment over intervals consistent with organizational risk monitoring cycles, and cross-cultural measurement equivalence testing across organizational contexts representing diverse Hofstede cultural dimension profiles, are identified as essential validation components before the HCRI can be applied with confidence in high-stakes governance decisions. This sequenced validation agenda reflects the epistemological standards of the behavioral measurement literature and positions the CHRF appropriately as a theoretically grounded framework in need of, rather than having completed, the empirical validation process that would justify its application as a precision risk management instrument.

2.2. Literature Search Protocol and Source Selection

To support transparency and auditability of the theoretical synthesis process, this subsection documents the literature search protocol, keyword strategy, and source selection criteria applied during framework development. Database Coverage. Literature was identified through searches of the following databases: PsycINFO, ACM Digital Library, IEEE Xplore, ProQuest Dissertations & Theses, Web of Science, and Google Scholar. Supplementary searches were conducted of the NIST National Cybersecurity Center of Excellence publication repository and the CISA resource library for governance-relevant documents. Search Terms. Core keyword clusters used across databases included: ("cyberpsychology" OR "cyber psychology") AND ("security behavior" OR "cyber risk" OR "human factors"); ("cognitive vulnerability" OR "cognitive bias") AND ("cybersecurity" OR "phishing" OR "social engineering"); ("behavioral susceptibility" OR "security compliance") AND ("information security" OR "insider threat"); ("affective influence" OR "emotional state" OR "occupational stress") AND ("security decision" OR "phishing susceptibility"); ("social engineering" OR "influence tactics") AND ("organi-

zational vulnerability” OR “security awareness”); (“security culture” OR “organizational context”) AND (“cybersecurity” OR “information security policy”). Foundational theoretical works, including Kahneman (2011), Cialdini (2007), Schein (2010), Bronfenbrenner (1979), and Suler (2004), were retrieved by author name and supplemented through backward and forward citation tracking from high-relevance empirical sources. Inclusion Criteria. Sources were included if they met at least one of the following criteria: 1) the source provides a theoretical mechanism directly informing one or more CHRF domain constructs; 2) the source reports peer-reviewed empirical findings documenting a measurable relationship between a psychological construct and a security-relevant behavioral outcome; 3) the source provides a validated psychometric instrument directly applicable to CHRF domain measurement; or 4) the source represents an authoritative governance or standards document from a recognized standards body (NIST, ISO, CISA, ENISA) with direct relevance to human risk assessment or behavioral security requirements.

Exclusion Criteria. Sources were excluded if they addressed cybersecurity from exclusively technical or network security perspectives without behavioral or psychological content; if they were practitioner or trade literature without peer review for empirical claims; or if their primary application domain was unrelated to organizational or individual security behavior. Source Retention. Across all databases and search rounds, an initial corpus of approximately 180 sources was identified through title and abstract review. Following full-text screening against the inclusion and exclusion criteria above, 42 sources were retained as primary references contributing directly to domain construction, measurement schema development, or governance alignment analysis. This article’s reference list reflects these retained primary sources. The review is characterized as a structured narrative synthesis rather than a systematic review because no formal search protocol was pre-registered, no PRISMA-style screening log was maintained, and source selection was guided by theoretical relevance judgment rather than algorithmic eligibility scoring.

3. Theoretical Foundations and Literature Review

3.1. Human Factors in Cybersecurity

The recognition that human behavior constitutes a primary attack surface in cybersecurity is not new, but the theoretical frameworks available to systematically study and mitigate this risk remain underdeveloped relative to the scope of the problem. [8] produced one of the earliest empirical investigations of user behavior in organizational security contexts, identifying that security policies frequently conflicted with work practices, leading to predictable behavioral workarounds. Subsequent scholarship by [9] formalized the concept of security behavior into distinct taxonomies of intentional and unintentional risk-generating actions, providing a behavioral science foundation upon which later models built.

The field of human factors engineering has contributed ergonomic and cognitive load perspectives to cybersecurity risk analysis [10]. Generic Error Modeling System and subsequent Human Error theories offered frameworks for understanding why individuals make predictable, systematic mistakes in complex technological environments. These cognitive error models proved especially relevant in the context of phishing susceptibility, password management failures, and misconfigured system administration, areas where human error intersects directly with security outcomes.

3.2. Cyberpsychology and Digital Behavior

Cyberpsychology as a scientific discipline emerged from the intersection of cognitive psychology, social psychology, and human-computer interaction research. [6] foundational work on the online disinhibition effect demonstrated that digital environments reduce social constraints in ways that alter both communication behavior and judgment. The reduction of normative cues, anonymity affordances, and asynchronous interaction patterns combine to produce behavioral profiles substantially different from offline contexts, a difference with profound implications for security decision-making.

[5] extended these insights by examining how social influence processes operate in computer-mediated communication, establishing that compliance, conformity, and persuasion dynamics are preserved and in some contexts amplified in digital environments. This finding directly informs the social engineering vulnerability domain of the CHRF, as it provides a psychological mechanism for understanding why phishing, vishing, and pretexting attacks achieve their documented success rates despite widespread awareness of these threats.

More recent cyberpsychology scholarship has examined the role of trust, identity, and personality in shaping online risk behavior. [7] demonstrated that individuals exhibit systematic differences in online fraud susceptibility linked to personality dimensions including conscientiousness, neuroticism, and openness to experience. [11] applied the elaboration likelihood model to phishing susceptibility, finding that individuals processing email under conditions of high cognitive load or emotional arousal were significantly more likely to comply with fraudulent requests.

3.3. Cognitive and Behavioral Frameworks

The CHRF draws explicitly upon [12] Dual-Process Theory of cognition, which distinguishes between fast, automatic, heuristic-driven System 1 thinking and slow, deliberate, analytical System 2 reasoning. Social engineering attacks are specifically designed to exploit System 1 processes by creating conditions of urgency, authority, scarcity, and emotional arousal that bypass deliberate evaluation. This cognitive architecture insight undergirds the Cognitive Vulnerability and Affective Influence domains of the CHRF and suggests that interventions targeting System 2 activation may reduce susceptibility to manipulation.

Behavioral economics contributions from [13] on choice architecture and nudge theory have been applied to cybersecurity contexts with promising results. [14] examined how behavioral nudges embedded in security interfaces could reduce risky behavior without requiring deliberate decision engagement, supporting the CHRF's emphasis on environmental and design-level interventions alongside individual-level training. Prospect Theory [15] further illuminates why users systematically underestimate cybersecurity risks when presented as probabilistic rather than experiential threats.

3.4. Organizational and Social Dimensions

The organizational dimension of human cyber risk has been examined through the lens of security culture, the set of values, beliefs, and behavioral norms within an organization that shape its members' security-relevant behaviors [16]. [17] developed one of the first empirically validated models of information security culture, demonstrating that organizational norms exert stronger predictive influence on security behavior than individual training or policy mandates. [18] extended this work by examining how managerial security commitment functions as a social proof mechanism that legitimizes, or delegitimizes, security-conscious behavior among employees.

Social influence processes documented in [19] foundational work on persuasion principles, reciprocity, commitment, social proof, authority, liking, and scarcity, provide the theoretical machinery through which adversaries construct social engineering attacks. The CHRF's Social Engineering Exposure domain operationalizes these principles as measurable vulnerability dimensions, enabling organizations to assess their workforce's differential susceptibility to specific influence tactics.

3.5. Existing Risk Frameworks and Their Limitations

Existing cybersecurity risk frameworks including NIST SP 800-30 [20], FAIR (Factor Analysis of Information Risk), and the CISA Risk and Vulnerability Assessment methodology offer structured approaches to organizational risk quantification. However, these frameworks treat human behavior primarily as a categorical variable (insider threat, user error) rather than as a multi-dimensional psychological construct subject to systematic measurement and modification. The Human Aspects of Information Security Questionnaire represents a significant advance in psychometric measurement of security behavior, but it does not integrate into a comprehensive risk architecture that links individual measurement to organizational-level risk quantification and governance [21].

The CHRF addresses these limitations by providing a theoretically grounded, multi-domain model that bridges psychological measurement, organizational assessment, and cybersecurity governance. It draws upon the strengths of existing frameworks while introducing the cyberpsychology depth necessary to capture the dynamic, contextual, and socially mediated nature of human cyber risk.

4. The Cyber Human Risk Framework

4.1. Architecture and Domain Exposition Framework Overview

The CHRF is structured around the foundational premise that human cyber risk is not a fixed trait of individuals but a dynamic, contextually modulated construct that emerges from the interaction of cognitive architecture, behavioral dispositions, emotional states, social influence exposure, and organizational context. The CHRF organizes these interacting variables into five primary domains: 1) Cognitive Vulnerability (CV), 2) Behavioral Susceptibility (BS), 3) Affective Influence (AI), 4) Social Engineering Exposure (SEE), and 5) Organizational Context (OC). These domains interact in non-linear, bidirectional ways, and their combined expression produces an individual or organizational Human Cyber Risk Index (HCRI) that enables risk stratification, intervention targeting, and longitudinal monitoring.

The CHRF adopts a systems perspective consistent with [22] ecological model of human development, recognizing that individual behavior is embedded within concentric layers of organizational, social, and technological context. Risk, in the CHRF model, is not localized in the individual alone but in the interaction between individual characteristics and environmental conditions, a perspective that implicates both workforce development and security environment design as legitimate intervention targets.

4.2. Domain I: Cognitive Vulnerability (CV)

Before elaborating each domain, a brief boundary statement is warranted for the three domains most susceptible to conceptual overlap: Cognitive Vulnerability (CV), Behavioral Susceptibility (BS), and Social Engineering Exposure (SEE). Cognitive Vulnerability is exclusively an internal psychological domain: it encompasses the cognitive processing characteristics, heuristic tendencies, mental model limitations, and attentional constraints that exist within the individual's cognitive architecture prior to any specific threat event or behavioral response. CV captures what the person is prone to cognitively; it does not include any observable behavior, any adversarial stimulus, or any characteristic of the organizational environment. Behavioral Susceptibility begins where CV ends it is exclusively an observable behavioral domain, capturing security-relevant actions and omissions that can be measured through performance data, system logs, simulation exercises, or structured self-report. BS includes what the person does or fails to do in security contexts; it does not include the cognitive mechanisms that generate those behaviors (which belong to CV) nor the external adversarial context that elicits them (which belongs to SEE). Social Engineering Exposure begins where both CV and BS end: it is exclusively an adversarial interface domain, capturing the characteristics of the individual's exposure to external social influence attempts, the digital footprint that adversaries can exploit, the influence tactic susceptibility that determines response to adversarial stimuli, and the organizational role that determines adversarial target priority. SEE does not include internal cognitive processing

(CV) or behavioral patterns independent of adversarial elicitation (BS). This boundary structure ensures that a phishing simulation click-through rate is classified as BS (observable behavior under standardized conditions), while an individual's susceptibility to authority-based urgency framing is classified as SEE (adversarial influence interface), and the heuristic processing tendency that underlies both is classified as CV (internal cognitive architecture). The three domains are causally connected, CV mechanisms predict BS outcomes, and SEE conditions determine which CV vulnerabilities adversaries target, but they are analytically and measurement-methodologically distinct.

Cognitive Vulnerability encompasses the range of cognitive characteristics, processing tendencies, and mental model limitations that predispose individuals to security-relevant errors. Drawing on cognitive psychology's extensive documentation of systematic heuristics and biases, the CHRF identifies four primary cognitive vulnerability sub-dimensions: attentional depletion, heuristic reliance, mental model inadequacy, and information overload susceptibility [11] [23].

Attentional depletion refers to the diminished capacity for security-relevant vigilance that accumulates through sustained cognitive work, decision fatigue, and competing task demands. Empirical research in occupational psychology has established that cognitive resources are finite and deplete across the workday, with implications for the timing and frequency of security-relevant decisions [24]. Security decisions requiring careful evaluation, reviewing email header information, scrutinizing URL legitimacy, verifying sender identity, are most vulnerable to attentional failure precisely when cognitive resources are lowest, typically in late afternoon or during high-workload periods.

Heuristic reliance captures the systematic tendency to substitute computationally efficient rules of thumb for analytically demanding evaluation. While heuristics are adaptive in the vast majority of everyday decisions, they create exploitable predictabilities in security contexts. Confirmation bias leads users to interpret ambiguous email communications in alignment with established expectations; authority bias drives compliance with requests attributed to figures of organizational power; availability heuristic causes individuals to underestimate the likelihood of threats they have not personally experienced.

Mental model inadequacy refers to gaps or misconceptions in an individual's conceptual representation of how cybersecurity threats operate. [25] documented through qualitative research that many users hold folk models of computer security that are systematically inaccurate, for example, believing that security threats are targeted rather than automated and indiscriminate, or that obvious visual signals distinguish legitimate communications from fraudulent ones. These mental model errors produce predictable failure patterns in security decision-making that training interventions must specifically target.

Information overload susceptibility captures individual differences in the ability to maintain security vigilance under conditions of high information volume and complexity. Modern digital work environments expose individuals to hun-

dreds of emails, notifications, system alerts, and communication requests daily, creating conditions that systematically degrade the deliberative processing required for security-relevant discrimination tasks. [26] documented the organizational consequences of information overload, and their findings apply directly to security decision quality in high-volume communication environments.

4.3. Domain II: Behavioral Susceptibility (BS)

Behavioral Susceptibility encompasses the observable security-relevant behaviors, and patterns of behavioral omission, that create exposure to cyber threats. Unlike Cognitive Vulnerability, which addresses internal processing tendencies, Behavioral Susceptibility is directly observable and measurable through self-report instruments, behavioral simulation exercises, and system log analysis. The CHRF identifies five primary behavioral susceptibility indicators: password hygiene compliance, phishing simulation performance, device and data handling practices, reporting behavior, and policy adherence fidelity.

Password hygiene compliance includes the cluster of behaviors surrounding credential management: password complexity, reuse patterns, storage practices, and multi-factor authentication adoption. Despite extensive documentation of password-related vulnerabilities and ubiquitous organizational policy requirements, empirical research consistently finds high rates of non-compliance [27]. The CHRF treats password hygiene as a behavioral rather than primarily knowledge-based variable, recognizing that awareness of best practices does not reliably translate into behavioral change without environmental support structures and motivational reinforcement.

Phishing simulation performance provides perhaps the most ecologically valid behavioral measure available for assessing susceptibility to social engineering attacks. Organizations that conduct simulated phishing exercises generate empirical behavioral data on individual and team-level susceptibility rates, click-through behaviors, and reporting propensity. The CHRF integrates phishing simulation metrics as a primary Behavioral Susceptibility indicator, normalized against industry benchmarks and adjusted for contextual factors including email content realism, organizational role, and prior simulation exposure.

Reporting behavior, the propensity to report suspected security incidents, policy violations, or social engineering attempts, functions as both a protective behavior indicator and an organizational security culture metric. [28] documented the significant gap between the frequency of potential security incidents and the frequency of formal reporting, attributing this gap to fear of blame, uncertainty about reporting thresholds, perceived futility of reporting, and organizational cultural factors that discourage admissions of error. The CHRF treats reporting behavior as a critical outcome variable that reflects both individual behavioral disposition and organizational contextual conditions.

4.4. Domain III: Affective Influence (AI)

Affective Influence encompasses the role of emotional states, stress responses, and

mood conditions in modulating security-relevant cognition and behavior. The CHRF recognizes that cybersecurity decision-making does not occur in an emotionally neutral vacuum; rather, emotional conditions systematically alter risk perception, deliberation quality, and behavioral compliance in ways that either amplify or attenuate vulnerability.

Stress represents the most extensively documented affective influence on security behavior. Research in occupational health psychology has established that acute and chronic stress impairs working memory capacity, narrows attentional focus, and increases reliance on heuristic processing [29]. In security contexts, stressed individuals exhibit elevated phishing susceptibility, increased rates of configuration errors, and reduced reporting behavior, all of which translate directly into elevated organizational cyber risk. The CHRF operationalizes occupational stress as a modulating variable that amplifies the expression of Cognitive Vulnerability and Behavioral Susceptibility.

Fear appeals, communications that invoke threat-based emotional responses to motivate security behavior change, represent a contested terrain in behavioral cybersecurity. While fear appeals have demonstrated efficacy in health communication research under specific conditions, their application in cybersecurity training contexts has produced inconsistent results [30]. Excessive fear arousal can produce defensive avoidance, denial, and disengagement rather than behavior change, particularly when individuals perceive low efficacy in their ability to respond effectively to the threat. The CHRF's Affective Influence domain incorporates efficacy beliefs as a critical moderating construct, consistent with Protection Motivation Theory [31].

Emotional manipulation is the intentional exploitation of affective states in social engineering attacks. Urgency induction, authority-based anxiety, reciprocity pressure, and sympathy arousal are all documented adversarial techniques for degrading deliberative processing and eliciting behavioral compliance. The CHRF's Affective Influence domain specifically measures individuals' susceptibility to emotion-modulated decision-making, providing an assessment dimension that connects social engineering attack methodology to individual psychological vulnerability profiles [32].

4.5. Domain IV: Social Engineering Exposure (SEE)

Social Engineering Exposure captures the intersection of adversarial social influence techniques and individual and organizational vulnerability to those techniques. Social engineering, the use of psychological manipulation to obtain unauthorized information or access, exploits the full range of human cognitive, behavioral, and affective vulnerabilities documented in the preceding domains. The CHRF's SEE domain addresses three interrelated exposure dimensions: influence tactic susceptibility, digital identity exposure, and open-source intelligence (OSINT) attack surface.

Influence tactic susceptibility assesses the degree to which individuals demon-

strate differential vulnerability to documented persuasion principles [19]. Research by [33] demonstrated that individuals with high scores on measures of social influence susceptibility, including tendencies toward authority deference, social conformity, and reciprocity responsiveness, exhibit significantly elevated phishing susceptibility rates. The CHRF incorporates validated psychometric instruments for influence susceptibility assessment, enabling risk stratification at the individual and role-specific level.

Digital identity exposure quantifies the degree to which an individual's personal and professional information is publicly accessible through social media platforms, professional networking sites, data broker repositories, and other open sources. Adversaries routinely conduct OSINT reconnaissance prior to targeted spear-phishing, pretexting, and whaling attacks, leveraging publicly available information to construct contextually credible lures. The CHRF's SEE domain includes a structured OSINT exposure assessment module that generates an individual Digital Footprint Risk Score (DFRS) based on the volume, sensitivity, and accessibility of an individual's publicly available information.

The organizational attack surface dimension of SEE extends beyond individual exposure to assess the collective social engineering vulnerability presented by an organization's workforce. High-risk roles, executive assistants, finance personnel, IT administrators, legal staff, present elevated social engineering targets due to their access privileges and procedural authority. The CHRF maps organizational role profiles to documented social engineering attack typologies, enabling security teams to prioritize protective interventions based on adversarial target value rather than applying uniform training approaches across all personnel.

4.6. Domain V: Organizational Context (OC)

Organizational Context encompasses the structural, cultural, leadership, and governance dimensions that shape the security behavior of all individuals within the organization. The CHRF positions Organizational Context as both a direct risk factor and a powerful moderator of all four individual-level domains, recognizing that individual psychological vulnerabilities express themselves differently, and with different organizational risk consequences, depending on the organizational environment in which they occur.

Security culture, the organization's shared values, assumptions, and norms regarding information security, functions as a foundational determinant of behavioral security outcomes. [17] demonstrated that organizations with positive security cultures exhibit lower rates of security incidents, higher reporting behavior, and more consistent policy compliance, independent of technical control quality. The CHRF incorporates security culture assessment through validated instruments including the Security Culture Framework, enabling organizations to benchmark their cultural security maturity against comparable organizations and identify specific cultural dimensions requiring intervention [33].

Leadership security commitment functions as a social modeling mechanism

that establishes the behavioral norms experienced by all personnel. When senior leadership visibly prioritizes security practices, accepting authentication inconveniences, participating in training exercises, supporting reporting without blame, this behavior signals organizational values that employees internalize and emulate. Conversely, leadership that openly circumvents security controls sends equally powerful behavioral signals that undermine policy-level security mandates. The CHRF assesses leadership security behavior as a distinct Organizational Context indicator through peer and subordinate assessment mechanisms.

Governance and policy quality assesses the completeness, usability, and enforcement consistency of an organization's security policies. Research by [34] found that perceived policy usefulness and top management support were the strongest predictors of employee information security policy compliance, outperforming organizational sanctions and monitoring as motivational drivers. The CHRF's governance assessment module evaluates policy characteristics against both legal compliance requirements and behavioral science principles of policy design that promote voluntary compliance.

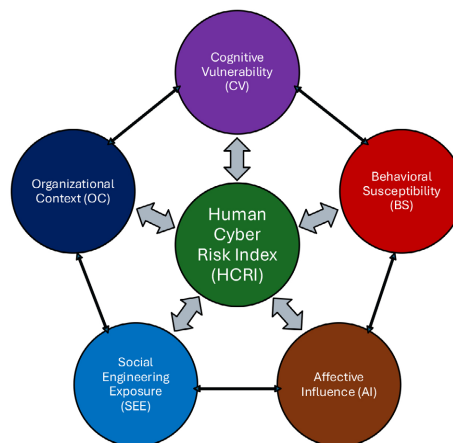


Figure 1. Cyber human risk framework.

Figure 1 the CHRF represents a systems-based model of human-driven cyber risk, where risk is not attributed to a single factor but emerges from the interaction of multiple psychological, behavioral, and organizational domains.

At the core of the figure is the Human Cyber Risk Index (HCRI), which serves as the central output metric. Surrounding this central node are five primary domains, Cognitive Vulnerability (CV), Behavioral Susceptibility (BS), Affective Influence (AI), Social Engineering Exposure (SEE), and Organizational Context (OC), each contributing to overall risk. The connecting arrows illustrate that these domains are interdependent and bidirectional, meaning changes in one domain can amplify or mitigate effects in others. For example, high stress (AI) can increase heuristic decision-making (CV), which may lead to riskier behaviors (BS).

The figure also reflects a nonlinear, ecological perspective, where individual risk is shaped by both internal characteristics and external environmental conditions,

particularly organizational culture and governance. The combined influence of all domains produces the HCRI score, which can then be used for risk stratification, targeted interventions, and continuous monitoring. Ultimately, the CHRF graph emphasizes that cybersecurity risk is a dynamic system-level phenomenon, not simply an individual user deficiency.

5. The Human Cyber Risk Index (HCRI)

5.1. Index Architecture and Rationale

The Human Cyber Risk Index (HCRI) is the CHRF's quantitative output mechanism, designed to translate multi-domain assessment data into a composite risk score that enables individual risk stratification, organizational benchmarking, and longitudinal risk monitoring. The HCRI is calculated as a weighted composite of domain-level scores, with weights derived from empirical research on the relative contribution of each domain to documented security incident outcomes.

The HCRI formula at the individual level is expressed as: $HCRI = w1 (CV) + w2 (BS) + w3 (AI) + w4 (SEE) + w5 (OC)$, where each domain score represents a normalized measure (0 - 100 scale) and the weights ($w1 - w5$) sum to 1.0. In the initial CHRF formulation, informed by the relative effect sizes documented in the behavioral cybersecurity literature, weights are provisionally set as: Cognitive Vulnerability (0.20), Behavioral Susceptibility (0.25), Affective Influence (0.15), Social Engineering Exposure (0.25), and Organizational Context (0.15). These weights should be treated as empirical hypotheses subject to validation and revision based on longitudinal outcome data.

The provisional HCRI weights, risk tier cut-points, and role-risk multipliers are not arbitrary; each is anchored to documented evidence and established conventions in the behavioral security and risk management literatures. The following justifications make these anchors explicit. **Domain Weights.** Behavioral Susceptibility and Social Engineering Exposure are each assigned the highest weight (0.25) for two convergent reasons. First, BS is the domain with the strongest direct empirical relationship to security incident outcomes: phishing simulation click-through rates, a primary BS indicator, are the single most widely replicated behavioral predictor of real-world phishing susceptibility documented in the organizational security literature [1] [2]. Second, SEE is assigned equivalent weight because adversarial social engineering is the most prevalent documented pathway to initial compromise in APT and criminal cyber operations, and the Digital Footprint Risk Score provides a measurable, reducible external risk driver that warrants equivalent analytical priority alongside behavioral outcomes [1]. Cognitive Vulnerability is assigned 0.20, reflecting its theoretical primacy as the mechanism through which adversarial social engineering operates, while acknowledging that CV's relationship to incident outcomes is partially mediated through BS and SEE rather than direct. The lower weights for Affective Influence (0.15) and Organizational Context (0.15) reflect their role as modulatory rather than primary drivers: both are well-documented amplifiers of CV and BS expression, but the behavioral

security literature documents smaller and more context-dependent direct effect sizes for affective and organizational variables than for behavioral and social engineering indicators, justifying their secondary weighting pending validation data that could support upward revision.

Risk Tier Cut-Points. The four-tier stratification schema, Low (0 - 24), Moderate (25 - 49), Elevated (50 - 74), and Critical (75 - 100), follows the quartile-based risk stratification convention established in NIST SP 800-30's risk level matrix and widely adopted in public health behavioral risk stratification (where four-tier schemas are standard in validated instruments including the PHQ-9 and PCL-5 severity classifications). Quartile-based cut-points on normalized 0 - 100 composite scales are methodologically appropriate when empirical outcome-based cut-points are unavailable during initial framework development, as they distribute the risk population across tiers without artificially concentrating cases at any threshold. The cut-points are explicitly provisional: once criterion validity testing links HCRI scores to documented incident outcomes, empirically derived cut-points should replace the current quartile conventions if the data support different threshold placements.

Role-Risk Multipliers. Role-risk multipliers in the RAO-HCRI are derived from three documented dimensions of differential organizational exposure: 1) access privilege level, grounded in the role-based access control literature's documentation of privilege escalation as the dominant pathway in insider threat and APT lateral movement; 2) organizational decision-making authority, reflecting the documented adversarial preference for targeting individuals who can authorize financial transactions, access sensitive data, or bypass normal approval workflows (a pattern documented extensively in business email compromise and whaling attack taxonomies); and 3) empirical adversarial targeting frequency by role category, as documented in Verizon DBIR sector-specific targeting data and social engineering attack taxonomy literature [1].

Executive and privileged-access roles are assigned the highest multipliers (suggested range: 1.5 - 2.0 $\times$) because they concentrate all three exposure dimensions; individual contributors with standard access and no authorization authority are assigned base multipliers (1.0 $\times$). Intermediate roles, team leads, finance staff, executive assistants, IT administrators, are assigned intermediate multipliers (suggested range: 1.2 - 1.5 $\times$) reflecting their elevated but not maximized access and targeting profiles. All multiplier values should be validated against the specific organizational threat environment and adjusted based on incident data if available.

Individual HCRI scores are classified into four risk tiers: Low (0 - 24), Moderate (25 - 49), Elevated (50 - 74), and Critical (75 - 100). Tier classification triggers differentiated intervention protocols, ranging from self-paced awareness reinforcement for Low-tier individuals to intensive behavioral coaching, role reassignment consideration, and enhanced monitoring for Critical-tier personnel. The tiered intervention model is consistent with public health risk stratification approaches and supports resource allocation efficiency by concentrating intensive interventions where risk is the highest.

Weight, Tier, and Multiplier Justification

The provisional HCRI weights, risk tier cut-points, and role-risk multipliers are not arbitrary; each is anchored to documented evidence and established conventions in the behavioral security and risk management literatures. The following justifications make these anchors explicit. Domain Weights. Behavioral Susceptibility and Social Engineering Exposure are each assigned the highest weight (0.25) for two convergent reasons. First, BS is the domain with the strongest direct empirical relationship to security incident outcomes: phishing simulation click-through rates, a primary BS indicator, are the single most widely replicated behavioral predictor of real-world phishing susceptibility documented in the organizational security literature [1] [2]. Second, SEE is assigned equivalent weight because adversarial social engineering is the most prevalent documented pathway to initial compromise in APT and criminal cyber operations [1], and the Digital Footprint Risk Score provides a measurable, reducible external risk driver that warrants equivalent analytical priority alongside behavioral outcomes. Cognitive Vulnerability is assigned 0.20, reflecting its theoretical primacy as the mechanism through which adversarial social engineering operates, while acknowledging that CV's relationship to incident outcomes is partially mediated through BS and SEE rather than direct.

The lower weights for Affective Influence (0.15) and Organizational Context (0.15) reflect their role as modulatory rather than primary drivers: both are well-documented amplifiers of CV and BS expression (28, 15, 16), but the behavioral security literature documents smaller and more context-dependent direct effect sizes for affective and organizational variables than for behavioral and social engineering indicators, justifying their secondary weighting pending validation data that could support upward revision. Risk Tier Cut-Points. The four-tier stratification schema, Low (0 - 24), Moderate (25 - 49), Elevated (50 - 74), and Critical (75 - 100), follows the quartile-based risk stratification convention established in NIST SP 800-30's risk level matrix and widely adopted in public health behavioral risk stratification (where four-tier schemas are standard in validated instruments including the PHQ-9 and PCL-5 severity classifications). Quartile-based cut-points on normalized 0 - 100 composite scales are methodologically appropriate when empirical outcome-based cut-points are unavailable during initial framework development, as they distribute the risk population across tiers without artificially concentrating cases at any threshold. The cut-points are explicitly provisional: once criterion validity testing links HCRI scores to documented incident outcomes, empirically derived cut-points should replace the current quartile conventions if the data support different threshold placements.

Role-Risk Multipliers. Role-risk multipliers in the RAO-HCRI are derived from three documented dimensions of differential organizational exposure: 1) access privilege level, grounded in the role-based access control literature's documentation of privilege escalation as the dominant pathway in insider threat and APT lateral movement [Cappelli *et al.*, 2012; CERT Insider Threat Center]; 2) organi-

zational decision-making authority, reflecting the documented adversarial preference for targeting individuals who can authorize financial transactions, access sensitive data, or bypass normal approval workflows (a pattern documented extensively in business email compromise and whaling attack taxonomies); and 3) empirical adversarial targeting frequency by role category, as documented in Verizon DBIR sector-specific targeting data and social engineering attack taxonomy literature [1] [32].

5.2. Organizational HCRI Aggregation

Organizational HCRI scores are computed through weighted aggregation of individual scores, with role-based weighting applied to account for the differential security impact of high-risk positions. An executive with Critical-tier HCRI scores represents a categorically different organizational risk than an individual contributor with equivalent scores, due to the executive’s elevated access privileges, decision-making authority, and adversarial target value. The CHRF’s organizational aggregation methodology applies role-risk multipliers derived from documented attack targeting patterns to produce a Role-Adjusted Organizational HCRI (RAO-HCRI) that accurately reflects the organization’s actual security risk exposure.

The organizational HCRI enables several high-value security governance applications. First, it provides a data-driven basis for security training investment prioritization, directing resources toward the highest-risk personnel and organizational units rather than applying uniform training approaches across all staff. Second, it generates a longitudinal risk trajectory that enables security leaders to demonstrate the return on investment of behavioral security interventions to executive stakeholders, a persistent challenge in cybersecurity governance that the HCRI directly addresses. Third, it enables cross-organization benchmarking when implemented across industry sectors, creating empirical norms that support sector-wide security culture advancement.

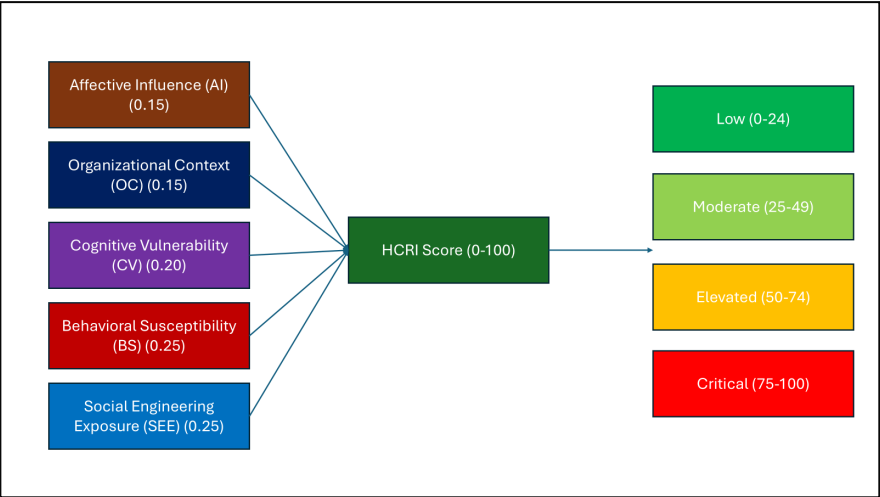


Figure 2. The human cyber risk index.

Figure 2 shows the HCRI functions as a multi-domain weighted composite index, where five psychologically and behaviorally grounded domains are normalized and aggregated into a single quantitative risk score. The model is additive in structure but interpretive in context, meaning identical scores may represent different risk realities depending on domain composition. The inclusion of role-based weighting at the organizational level transforms the model from a simple average into a risk-sensitive aggregation system, aligning cybersecurity decision-making with real-world adversarial targeting patterns.

5.3. Domain Operationalization and Scoring Protocol

For the HCRI formula to generate interpretable and comparable scores, each domain must be operationalized through consistent data sources, measurement instruments, and normalization rules. The following specifications define the intended operationalization for each domain in the initial CHRF implementation. These specifications are provisional and should be refined through the validation sequence described in Section 8.1; they are provided here to make the scoring architecture auditable and to guide practitioners implementing the framework prior to validation. Cognitive Vulnerability (CV). CV is operationalized through a composite of three data sources: 1) a validated psychometric instrument assessing heuristic reliance and cognitive bias susceptibility (recommended: the Cognitive Reflection Test [CRT] for analytical reasoning propensity, supplemented by the Need for Cognition Scale for deliberative processing tendency); 2) mental model accuracy assessment administered through scenario-based security judgment tasks that compare participant threat assessments against expert ground truth; and 3) attentional depletion and information overload indicators derived from cognitive load self-report scales adapted for organizational security contexts. Raw scores on each sub-instrument are normalized to a 0 - 100 scale using z-score standardization against role-comparable organizational norms, then averaged to produce the composite CV score. Higher scores indicate greater cognitive vulnerability.

Behavioral Susceptibility (BS). BS is operationalized through a composite of directly observable behavioral indicators: 1) phishing simulation click-through and credential submission rates normalized against industry benchmarks by sector and role type, expressed as percentile-referenced scores; 2) password hygiene compliance scores derived from system log or self-report audit of credential management behaviors (complexity, reuse, MFA adoption), scored against policy-defined compliance thresholds; 3) policy adherence fidelity rates from security policy compliance audits; and 4) incident reporting propensity, operationalized as the ratio of self-reported near-miss events to formally filed incident reports over a defined monitoring period. Each behavioral indicator is normalized to 0 - 100 (higher = greater susceptibility), then composited with equal sub-weights within the BS domain pending validation evidence for differential weighting. Affective Influence (AI). AI is operationalized through: 1) an occupational stress measure appropriate to the organizational context (recommended: the Perceived Stress

Scale [PSS] adapted for workplace security contexts, or the Job Stress Scale where occupational norms are available); 2) an efficacy belief assessment measuring perceived response efficacy and self-efficacy in security threat response, consistent with Protection Motivation Theory operationalizations; and 3) an emotional manipulation susceptibility scale assessing differential vulnerability to urgency, authority, and reciprocity-based emotional appeals in security decision scenarios. Raw scores are standardized and aggregated to produce the composite AI score (higher = greater affective influence on security behavior). Social Engineering Exposure (SEE). SEE is operationalized through: 1) a validated influence susceptibility instrument assessing differential vulnerability to documented persuasion principles (authority deference, social conformity, reciprocity responsiveness); 2) a structured Digital Footprint Risk Score (DFRS) derived from OSINT assessment of publicly accessible personal and professional information across social media, professional networking sites, and data broker repositories, scored on volume, sensitivity, and accessibility dimensions; and 3) a role-based attack surface score derived from mapping the individual's organizational role against documented adversarial targeting patterns (see Section 4.5).

The three SEE sub-scores are normalized to 0 - 100 and composited, with the DFRS and role-based attack surface sub-scores receiving supplementary weight where the individual's role presents elevated adversarial targeting risk. Organizational Context (OC). OC is operationalized through organizational-level rather than individual-level data sources: 1) a validated security culture assessment instrument (recommended: the Security Culture Framework assessment battery or equivalent, producing dimension scores for leadership commitment, peer behavioral norms, reporting culture, and policy legitimacy perception); 2) a governance and policy quality audit assessing policy completeness, usability, enforcement consistency, and alignment with behavioral science compliance principles; and 3) leadership security behavior ratings derived from multi-source assessment (peer and subordinate ratings of security-relevant leadership conduct). OC scores are computed at the organizational unit level and applied uniformly to all individuals within that unit, reflecting OC's role as a contextual moderator rather than an individual difference variable.

Missing Data Protocol. CHRF scoring is designed for full-domain operationalization, but practitioners may face incomplete data, particularly in resource-constrained environments. The following rules govern scoring when domain inputs are unavailable: 1) when fewer than two sub-indicators are available for any domain, that domain score must be flagged as 'indeterminate' rather than imputed, and the HCRI composite must be reported with an explicit data-completeness caveat; 2) when one sub-indicator is missing from a domain with three or more available sub-indicators, the composite may be calculated from available indicators with notation of the missing element; 3) when the full OC domain assessment is unavailable (common in small organizations), a conservative default OC score of 50 (mid-range) may be applied as a neutral assumption, with explicit flagging

of the imputation; and 4) HCRI scores derived from incomplete domain data must not be used for high-stakes individual personnel decisions without supplementary qualitative evidence. Practitioners should prioritize data completeness for the two highest-weighted domains, BS and SEE, as these contribute 50% of the composite HCRI weight and carry the greatest consequence for risk stratification accuracy.

6. Framework Integration with Cybersecurity Governance Standards

6.1. NIST Cybersecurity Framework Alignment

The CHRF is designed for integration with established cybersecurity governance standards rather than replacement of them. The NIST Cybersecurity Framework (CSF) 2.0 organizes cybersecurity activities around six core functions: Govern, Identify, Protect, Detect, Respond, and Recover. The CHRF's domain architecture maps across all six functions, with particular relevance to the Identify and Protect functions, which encompass asset management, risk assessment, awareness training, and security policy implementation [33].

Within the NIST CSF's Identify function, the CHRF's five-domain assessment protocol directly advances the organizational understanding of cybersecurity risk required by the ID.RA risk assessment category. The HCRI provides quantitative human risk data that can be integrated with technical vulnerability assessment data to generate a comprehensive, multi-dimensional organizational risk profile. This integration enables NIST CSF practitioners to move from categorical human risk acknowledgment ("users are the weakest link") to quantified, domain-specific human risk characterization that supports targeted risk treatment decisions [34].

6.2. ISO/IEC 27001 Integration

ISO/IEC 27001:2022 [3] requires organizations to assess information security risks to the confidentiality, integrity, and availability of information assets and to implement appropriate controls. The CHRF supports ISO 27001 compliance by providing a structured methodology for assessing the human risk dimension of information security, an area addressed by the standard's Annex A control domains relating to human resource security, access control, and information security awareness training.

The CHRF's Organizational Context domain assessment is particularly well-aligned with ISO 27001's requirements for leadership commitment, organizational role definitions, and security awareness programs. CHRF assessment outputs can serve as documented evidence of risk assessment due diligence required for ISO 27001 certification and ongoing compliance monitoring [35].

6.3. Integration with Insider Threat Programs

The CHRF has particular relevance for organizational insider threat programs, which seek to identify and mitigate risks posed by current or former employees, contractors, and business partners who have or had authorized access to organi-

zational assets. [33] established foundational profiles of insider threat actors that emphasize psychological and behavioral precursors, and subsequent research by the CERT Insider Threat Center documented the predictive validity of behavioral indicators in insider threat identification [36].

The CHRF's multi-domain assessment framework provides a structured, privacy-protective methodology for identifying individuals at elevated insider risk without the stigmatizing profiling approaches that have characterized earlier insider threat models. By focusing on behavioral security indicators and organizational context factors rather than psychological trait profiling, the CHRF supports insider threat risk management within appropriate civil liberties and workplace privacy boundaries.

7. Practical Applications and Implementation Guidance

7.1. Workforce Security Training and Development

The CHRF's most immediate practical application is in the design and targeting of workforce security training programs. Traditional security awareness training approaches rely on broad-coverage, one-size-fits-all curricula delivered at annual intervals, an approach that behavioral research has consistently shown to be ineffective at producing durable behavior change [37]. The CHRF enables a fundamentally different approach: precision security education targeting the specific cognitive, behavioral, and affective vulnerability dimensions measured for each individual or role category.

Individuals identified with elevated Cognitive Vulnerability scores benefit from training that explicitly addresses cognitive bias mechanisms and provides metacognitive strategies for recognizing when System 1 processing is operating in high-stakes security decision contexts. This approach draws on the educational psychology literature on debiasing training, which has demonstrated that explicit instruction in bias mechanisms, combined with practice in applying corrective strategies, produces measurable reductions in heuristic error rates that persist over time [38].

Individuals with elevated Behavioral Susceptibility scores, particularly on phishing simulation metrics, benefit from repeated, realistic behavioral practice in identifying social engineering attempts combined with immediate performance feedback. This approach is consistent with deliberate practice theory, which establishes that skill acquisition in complex, high-stakes domains require focused practice with expert feedback rather than passive information exposure [39]. Simulated phishing exercises, when designed with behavioral science principles, function as deliberate practice mechanisms for social engineering resistance skills.

7.2. Security Culture Development

The CHRF's Organizational Context domain assessment provides security culture practitioners with a multi-dimensional diagnostic that enables targeted culture development interventions. Rather than addressing security culture as a mono-

lithic construct, the CHRF's assessment reveals the specific cultural dimensions, leadership commitment, peer behavioral norms, reporting culture, policy legitimacy perception, that are most deficient and most impactful for a given organization.

[15] model of organizational culture change provides the theoretical foundation for CHRF-informed security culture development, emphasizing the primacy of leadership behavior, artifact creation, and espoused value alignment in shifting underlying cultural assumptions. Security culture development programs informed by the CHRF's diagnostic data can sequence interventions strategically, beginning with leadership behavior modeling, followed by policy redesign aligned with behavioral science principles, followed by peer norm reinforcement mechanisms, in a progression that addresses the deepest cultural layers rather than surface-level awareness campaigns.

7.3. Cyber Incident Response Enhancement

The CHRF has significant implications for cyber incident response by providing responders with a behavioral science framework for understanding how human factors contributed to incident occurrence and how they may affect response effectiveness. Post-incident HCRI analysis enables organizations to identify the specific cognitive, behavioral, or affective conditions that created the vulnerability exploited in an attack, information that is essential for designing targeted corrective interventions rather than generic remediation responses.

Furthermore, the high-stress conditions of active incident response create significant Affective Influence vulnerabilities among responders themselves, as decision-making quality degrades under the acute stress and time pressure of crisis conditions. The CHRF's Affective Influence domain provides a conceptual framework for incident response team leaders to build in stress mitigation practices, structured decision checklists, role separation, mandatory rest rotations, after-action emotional processing, that preserve cognitive function and behavioral security discipline throughout the response period.

8. Research Implications and Future Directions

8.1. Empirical Validation Requirements

The CHRF is introduced as a theoretically grounded framework requiring rigorous empirical validation before its quantitative elements should be applied with high confidence in high-stakes governance contexts. Validation priorities include: 1) psychometric validation of domain-specific assessment instruments across diverse organizational and demographic populations; 2) predictive validity testing of the HCRI against documented security incident outcomes in organizational settings; 3) test-retest reliability assessment of domain scores over time intervals relevant to longitudinal risk monitoring; and 4) cross-cultural validation of the framework's applicability across organizational contexts characterized by different cultural orientations toward authority, uncertainty, and individualism-collec-

tivism dimensions that may moderate the expression of CHRF domains [40].

Randomized controlled trial designs offer the most rigorous validation methodology for testing the CHRF's intervention prescriptions, comparing CHRF-guided precision security training against conventional broad-coverage approaches on behavioral outcome metrics. Given the ethical and logistical constraints of withholding security training from control groups, stepped-wedge and waitlist control designs may offer practical compromises that preserve internal validity while enabling ethical implementation.

8.2. Methodological Contributions to Cyberpsychology

The CHRF contributes to cyberpsychology methodology by articulating a multi-domain, multi-method assessment philosophy that draws on the discipline's characteristic integration of survey research, behavioral simulation, physiological measurement, and system log analysis. Future CHRF research should explore the integration of neuropsychological indicators, including pupillometry, galvanic skin response, and functional neuroimaging correlates of security decision-making, with behavioral and self-report measures to build a more complete model of the cognitive-affective architecture underlying human cyber risk.

Machine learning approaches offer promising avenues for HCRI refinement, enabling dynamic weight adjustment based on continuously accumulating outcome data and identification of non-linear interaction effects among CHRF domains that linear composite scoring cannot capture. [41] demonstrated the feasibility of applying predictive analytics to security behavior data, and subsequent advances in explainable AI (XAI) methods offer pathways for deploying machine learning-enhanced HCRI scoring without sacrificing the interpretability essential for practitioner adoption.

8.3. Policy Implications

The CHRF's development has significant policy implications for national cybersecurity workforce strategy. The Cybersecurity and Infrastructure Security Agency (CISA) National Cyber Workforce and Education Strategy identifies behavioral and human-centered cybersecurity competencies as priority development areas for the national workforce [42]. The CHRF provides the theoretical architecture necessary to operationalize these policy priorities into assessment instruments, curriculum frameworks, and competency standards that can be adopted across the public and private sectors.

At the organizational governance level, regulatory frameworks including HIPAA Security Rule requirements, NIST SP 800-171 controls for Controlled Unclassified Information protection, and emerging SEC cybersecurity disclosure requirements increasingly mandate documented human risk assessment and training effectiveness measurement. The CHRF's HCRI measurement architecture provides a structured, documentable approach to these compliance requirements that integrates scientific rigor with governance utility.

8.4. Limitations and Boundary Conditions

The CHRF, as an original theoretical framework, carries inherent limitations that practitioners and researchers must acknowledge. First, the provisional domain weights used in the HCRI calculation represent informed hypotheses derived from existing literature rather than empirically calibrated parameters derived from original data collection. Practitioners should treat HCRI scores as directional indicators rather than precise risk measurements until validation studies establish the metric properties of the composite index.

Second, the CHRF's development is grounded primarily in Western organizational and psychological research traditions. The applicability of framework assumptions, particularly in the Organizational Context domain, to organizations operating in high power-distance, uncertainty-avoidance, or collectivist cultural contexts requires specific validation. Psychological constructs including cognitive authority deference, reporting culture, and policy legitimacy perception carry different meanings and behavioral implications across cultural frameworks, and CHRF assessments deployed in non-Western organizational contexts should incorporate appropriate cultural adaptation.

Third, the CHRF does not currently address the full spectrum of insider threat motivation, focusing primarily on unintentional and manipulation-driven human risk rather than deliberate malicious insider behavior. While the Behavioral Susceptibility and Organizational Context domains provide partial coverage of disaffection indicators documented in insider threat research, the CHRF's psychometric approach is better suited to assessing vulnerability than intent. Organizations with significant insider threat concerns should supplement CHRF assessment with threat-specific behavioral indicator frameworks [35].

Fourth, the CHRF's measurement architecture assumes access to organizational data, including training records, phishing simulation results, incident report data, and policy compliance audits, that may not be available in all organizational contexts. Small and medium-sized enterprises, in particular, may lack the infrastructure to generate the behavioral data streams required for full HCRI operationalization. Abbreviated CHRF assessment protocols suitable for resource-constrained contexts represent a priority development need for broader framework adoption.

9. Discussion

The proliferation of sophisticated cyber threats has increasingly exposed a foundational inadequacy in how organizations conceptualize security risk: the human element is not a peripheral concern but the dominant attack surface. With over 74% of all data breaches consistently involving a human factor, through error, social engineering, credential misuse, or insider behavior, the field faces not a knowledge gap but a structural problem rooted in cognitive architecture, behavioral disposition, emotional state, and social vulnerability. Yet the dominant governance frameworks, NIST CSF, ISO/IEC 27001, CIS Controls, continue to treat

human behavior primarily as a compliance variable rather than a dynamic, scientifically tractable risk domain. The CHRF is constructed to address that gap directly, positioning cyberpsychology not as an academic interest but as an applied discipline with direct operational utility for risk assessment, workforce development, and organizational security culture.

What distinguishes the CHRF from prior human-factor models is the rigor of its interdisciplinary theoretical synthesis. Rather than drawing from a single tradition, the framework integrates cognitive psychology, behavioral economics, organizational behavior, social influence theory, and cyberpsychology into a coherent multi-domain architecture where each discipline addresses a distinct mechanism through which human vulnerability is created and exploited. Kahneman's Dual-Process Theory is the most foundational underpinning: the distinction between fast, heuristic-driven System 1 processing and deliberate, analytical System 2 reasoning maps directly onto adversarial social engineering methodology. Phishing campaigns, vishing attacks, and pretexting operations are specifically designed to exploit System 1 conditions, urgency, authority signaling, scarcity framing, emotional arousal, that bypass the deliberative evaluation in which security competencies reside. This insight is not merely descriptive but prescriptive: interventions that train individuals to recognize when System 1 is operating in high-stakes security contexts produce measurable reductions in susceptibility. Behavioral economics extends this through nudge theory and Prospect Theory, demonstrating that probabilistic threats are temporally discounted in ways that experiential demands are not, and that environmental design, choice architecture, system defaults, interface structure, can reduce risky behavior without requiring deliberate engagement. [18] social influence principles provide the theoretical machinery through which adversarial social engineering is constructed and through which individual susceptibility can be measured. Suler's work on the online disinhibition effect and Joinson's findings on compliance amplification in computer-mediated communication establish that digital environments do not merely replicate offline vulnerability, they amplify it. Organizational psychology, through Schein's cultural change model and Da Veiga and Elof's empirical security culture research, anchors the CHRF's recognition that individual behavior is embedded within and shaped by organizational norms, leadership behavior, and policy legitimacy perception in ways that make organizational context as much a risk variable as individual psychology.

The framework organizes human cyber risk into five primary domains, Cognitive Vulnerability, Behavioral Susceptibility, Affective Influence, Social Engineering Exposure, and Organizational Context, that interact bidirectionally and whose combined expression produces an individual or organizational Human Cyber Risk Index. This architecture reflects a deliberate commitment to precision over simplification: human cyber risk is not a fixed individual trait but a contextually modulated construct that emerges from the interaction of these domains, and understanding that interaction is essential for designing interventions that address

root causes rather than surface symptoms. Cognitive Vulnerability encompasses four sub-dimensions with direct security relevance. Attentional depletion, the finite cognitive resource degradation that accumulates through sustained work and decision fatigue, creates predictable vulnerability windows during late-day and high-workload periods precisely when security decisions requiring careful evaluation are most likely to occur. This is not merely a training problem but a work design problem, and the CHRF's articulation of it as a measurable domain creates the conceptual foundation for organizational interventions at the workflow and policy level. Mental model inadequacy captures the systematic inaccuracies in how individuals understand how threats operate, the belief that attacks are targeted rather than automated, or that visual authenticity cues reliably distinguish legitimate from fraudulent communications, and implies training approaches that explicitly surface and restructure inaccurate security schemas rather than layering accurate information atop unchallenged misconceptions. Heuristic reliance and information overload susceptibility complete the domain by documenting the predictable cognitive shortcuts and attentional limitations that create exploitable regularities in security decision-making under realistic organizational conditions.

Behavioral Susceptibility's emphasis on reporting behavior as a critical outcome variable is an underappreciated insight. The large gap between the frequency of potential security incidents and formal reporting, attributable to fear of blame, perceived futility, and organizational cultural factors, locates the solution not in individual training but in leadership behavior and organizational culture. This connection is captured through the bidirectional relationship the CHRF establishes between Behavioral Susceptibility and Organizational Context, reflecting the framework's systems perspective. Phishing simulation performance, when treated as a deliberate practice mechanism rather than merely an assessment tool, with realistic contextual variation, immediate feedback, and progressive difficulty, functions as a genuine skill development intervention consistent with Ericsson and colleagues' foundational research on expertise acquisition in complex domains.

The Affective Influence domain introduces a dimension that organizational cybersecurity frameworks have almost entirely neglected: emotional states as modulators of security-relevant cognition. Occupational stress impairs working memory, narrows attentional focus, and increases heuristic reliance in ways that translate directly into elevated phishing susceptibility, increased configuration error rates, and reduced reporting behavior. This creates a direct pathway from employee wellbeing to security incident risk, a link that justifies integrating occupational health considerations into security risk architecture. The CHRF's treatment of fear appeals in security communication reflects similar sophistication: excessive fear arousal produces defensive avoidance and disengagement rather than behavior change, particularly when individuals perceive low response efficacy. Protection Motivation Theory's requirement for joint high threat appraisal and high efficacy appraisal means that fear-based security communication without parallel efficacy-building content is not merely ineffective but potentially counterproductive, a nu-

ance that separates a psychologically grounded framework from a governance checklist.

Social Engineering Exposure's inclusion of a Digital Footprint Risk Score based on OSINT exposure assessment addresses a dimension of vulnerability that is measurable, reducible, and largely unaddressed by existing frameworks. Spear-phishing and whaling campaigns achieve elevated success rates precisely because they are contextually credible, that credibility derives from publicly accessible information that adversaries harvest in reconnaissance. Quantifying an individual's OSINT-accessible information as a component of their social engineering exposure risk creates a manageable variable addressable through digital hygiene practices and targeted education without requiring changes to underlying psychological characteristics. The domain's role-based targeting of organizational attack surface, mapping personnel roles to adversarial target typologies, addresses a systematic inefficiency in current security training allocation. Concentrating intensive intervention resources where adversarial pressure is actually the highest, rather than distributing training uniformly, represents a resource allocation improvement that the CHRF's analytical structure directly enables.

The Human Cyber Risk Index translates multi-domain assessment into a governance-ready instrument through a weighted composite formula with four-tier risk stratification that enables differentiated, resource-efficient intervention architecture. The Role-Adjusted Organizational HCRI, incorporating role-risk multipliers that account for access privileges and adversarial target value, captures the categorical difference in organizational exposure between a Critical-tier executive and a Critical-tier individual contributor, a distinction that flat aggregation approaches cannot represent. Three governance applications of the HCRI are particularly significant: data-driven training investment prioritization that allocates resources to highest-risk personnel rather than organizational hierarchy; longitudinal risk trajectory data that enables security leaders to demonstrate return on investment to executive stakeholders who currently receive few credible behavioral security metrics; and cross-organization benchmarking that creates the empirical norm infrastructure necessary for sector-wide security culture advancement. The CHRF's positioning as a complement to existing governance standards rather than a competing framework, with explicit alignment to NIST CSF 2.0 and ISO/IEC 27001, ensures that organizations need not abandon existing governance investments but can enhance them by adding the human risk assessment depth those frameworks presently lack.

The framework's practical implications for security culture development draw on Schein's organizational culture change model to articulate intervention sequencing logic that security programs most commonly miss. Leadership behavior modeling establishes cultural norms through social proof mechanisms that no training campaign can replicate; policy redesign aligned with behavioral science principles of voluntary compliance rather than purely punitive deterrence follows; peer norm reinforcement mechanisms sustain the shift at the cultural substrate

rather than at the level of surface awareness. The application of debiasing training research to Cognitive Vulnerability reduction offers a particularly powerful practical direction. Morewedge and colleagues' demonstration that explicit instruction in bias mechanisms combined with corrective strategy practice produces measurable and persistent reductions in heuristic error rates implies training approaches that teach individuals how their own cognitive architecture creates vulnerability, rather than merely teaching them what threats look like.

The CHRF's acknowledged limitations are important boundary conditions for practitioners. The provisional domain weights in the HCRI are empirical hypotheses requiring validation, not established parameters, and scores should be treated as directional indicators until psychometric validation studies establish their metric properties. The cross-cultural validity concern is substantial: constructs central to the framework, authority deference, reporting culture norms, policy legitimacy perception, carry different behavioral expressions across cultural frameworks characterized by different power distance, uncertainty avoidance, and individualism-collectivism orientations. Organizations operating across multinational workforces cannot assume that instruments validated in Western organizational contexts will generalize without cultural adaptation. The framework is also better suited to assessing vulnerability than deliberate malicious intent, meaning organizations with significant insider threat concerns will need to supplement CHRF assessment with threat-specific behavioral indicator frameworks. The infrastructure requirements for full HCRI operationalization, phishing simulation systems, training records, incident report data, policy compliance audits, create an equity gap that limits adoption to large enterprises while leaving resource-constrained organizations, where behavioral security programs are often most underdeveloped, without an accessible implementation pathway.

The urgency of the CHRF's contribution is amplified by the adversarial trajectory. AI-generated deepfakes, synthetic persona operations, and adaptive social engineering platforms capable of dynamically optimizing influence tactics based on behavioral response data each directly amplify the specific mechanisms the CHRF maps, exploiting authority heuristics with unprecedented contextual realism, constructing relationship histories that bypass authenticity checks, and targeting affective influence dimensions at the individual level with a precision that rule-based social engineering cannot achieve. Organizations that continue to rely on annual awareness training, categorical risk labels, and compliance-oriented governance responses will find those approaches increasingly inadequate as adversaries field capabilities specifically designed to exploit human cognitive architecture at scale. The CHRF offers the conceptual infrastructure for a response proportionate to that threat, scientifically grounded, multi-domain, quantifiable, and designed for empirical refinement as the validation research agenda it articulates accumulates evidence. Its most important contribution may ultimately be not the framework itself but the research program it invites: a systematic, interdisciplinary empirical effort to translate cyberpsychology's scientific depth into the pre-

cision risk management tools that the human dimension of cybersecurity has urgently needed and has not yet had.

10. Conclusions

The Cyber Human Risk Framework represents a significant advance in the theoretical and practical instrumentation of human-centered cybersecurity risk management. By synthesizing contributions from cognitive psychology, behavioral economics, organizational behavior, social influence theory, and cyberpsychology into a coherent, multi-domain architecture, the CHRF provides the field with a framework capable of bridging the persistent gap between scientific understanding of human behavior and the applied demands of cybersecurity governance.

The framework's five domains, Cognitive Vulnerability, Behavioral Susceptibility, Affective Influence, Social Engineering Exposure, and Organizational Context, collectively capture the multi-dimensional nature of human cyber risk with a fidelity that categorical approaches cannot achieve. The Human Cyber Risk Index translates this multi-domain assessment into a governance-ready quantitative instrument that enables risk stratification, intervention targeting, longitudinal monitoring, and regulatory compliance documentation.

The CHRF is explicitly positioned as a living framework requiring continuous empirical refinement as validation studies accumulate evidence on its psychometric properties, predictive validity, and cross-cultural applicability. The field of cyberpsychology is uniquely positioned to lead this validation agenda, drawing upon its distinctive methodological repertoire and its commitment to understanding human behavior in digital contexts with the rigor and precision that the stakes of contemporary cybersecurity demand.

Organizations that adopt the CHRF as a component of their cybersecurity risk management architecture gain access to a theoretically grounded, ethically defensible, and structured approach to human-centered risk assessment whose practical applicability will be strengthened as empirical validation studies establish its psychometric properties, predictive validity, and cross-cultural generalizability. In its current form, the CHRF provides a scientifically informed conceptual architecture and a provisional measurement schema; organizations should implement its quantitative elements with appropriate recognition of their pre-validation status and supplement HCRI scores with qualitative judgment, particularly for high-stakes personnel or governance decisions.

The CHRF is offered to the cyberpsychology and cybersecurity communities as a theoretical contribution, a practical instrument, and an invitation to collaborative empirical validation that will strengthen its scientific foundations and broaden its practical applications across the full spectrum of organizations navigating the complex terrain where human psychology meets digital risk.

Conflicts of Interest

The author declares no conflict of interest regarding the publication of this paper.

References

- [1] Dayama, R. (2025) Network and Infrastructure Security. *International Journal of Original Recent Advanced Research*, **2**.
<https://ijorarjournal.com/web/Download/2026-02-14-12-03-39-Per%20Id%200928.pdf>
- [2] Bada, M., Sasse, A.M. and Nurse, J.R.C. (2019) Cyber Security Awareness Campaigns: Why Do They Fail to Change Behaviour? arXiv: 1901.02672.
- [3] National Institute of Standards and Technology (2018) Framework for Improving Critical Infrastructure Cybersecurity (NIST Cybersecurity Framework, Version 1.1). U.S. Department of Commerce. <https://doi.org/10.6028/NIST.CSWP.04162018>
- [4] International Organization for Standardization (2022) ISO/IEC 27001:2022: Information Security, Cybersecurity and Privacy Protection, Information Security Management Systems, Requirements. ISO.
- [5] Joinson, A.N. (2007) Disinhibition and the Internet. In: Gackenbach, J., Ed., *Psychology and the Internet: Intrapersonal, Interpersonal, and Transpersonal Implications (2nd Edition)*, Academic Press, 75-92.
- [6] Suler, J. (2004) The Online Disinhibition Effect. *CyberPsychology & Behavior*, **7**, 321-326. <https://doi.org/10.1089/1094931041291295>
- [7] Whitty, M.T. and Buchanan, T. (2012) The Online Romance Scam: A Serious Cyber-crime. *Cyberpsychology, Behavior, and Social Networking*, **15**, 181-183.
<https://doi.org/10.1089/cyber.2011.0352>
- [8] Adams, A. and Sasse, M.A. (1999) Users Are Not the Enemy. *Communications of the ACM*, **42**, 40-46. <https://doi.org/10.1145/322796.322806>
- [9] Stanton, J.M., Stam, K.R., Mastrangelo, P. and Jolton, J. (2005) Analysis of End User Security Behaviors. *Computers & Security*, **24**, 124-133.
<https://doi.org/10.1016/j.cose.2004.07.001>
- [10] Reason, J. (1990) Human Error. Cambridge University Press.
<https://doi.org/10.1017/cbo9781139062367>
- [11] Vishwanath, A., Herath, T., Chen, R., Wang, J. and Rao, H.R. (2011) Why Do People Get Phished? Testing Individual Differences in Phishing Vulnerability within an Integrated, Information Processing Model. *Decision Support Systems*, **51**, 576-586.
<https://doi.org/10.1016/j.dss.2011.03.002>
- [12] Kahneman, D. (2011) Thinking, Fast and Slow. Farrar, Straus and Giroux.
- [13] Thaler, R.H. and Sunstein, C.R. (2008) Nudge: Improving Decisions about Health, Wealth, and Happiness. Yale University Press.
- [14] Renaud, K. and Zimmermann, V. (2019) Nudging Folks towards Stronger Password Choices: Providing Certainty Is the Key. *Behavioural Public Policy*, **3**, 228-258.
- [15] Kahneman, D. and Tversky, A. (1979) Prospect Theory: An Analysis of Decision under Risk. *Econometrica*, **47**, 263-292. <https://doi.org/10.2307/1914185>
- [16] Schein, E.H. (2010) Organizational Culture and Leadership. 4th Edition, Jossey-Bass.
- [17] Da Veiga, A. and Eloff, J.H.P. (2010) A Framework and Assessment Instrument for Information Security Culture. *Computers & Security*, **29**, 196-207.
<https://doi.org/10.1016/j.cose.2009.09.002>
- [18] Sommestad, T., Karlzén, H., Nilsson, P. and Hallberg, J. (2016) An Empirical Test of the Perceived Relationship between Risk and the Constituents Severity and Probability. *Information & Computer Security*, **24**, 194-204.
<https://doi.org/10.1108/ics-01-2016-0004>

- [19] Cialdini, R.B. (2007) *Influence: The Psychology of Persuasion*. Harper-Collins.
- [20] National Institute of Standards and Technology (2012) *Guide for Conducting Risk Assessments* (NIST SP 800-30, Rev. 1). U.S. Department of Commerce.
<https://doi.org/10.6028/NIST.SP.800-30r1>
- [21] Parsons, K., McCormac, A., Butavicius, M., Pattinson, M. and Jerram, C. (2014) Determining Employee Awareness Using the Human Aspects of Information Security Questionnaire (HAIS-Q). *Computers & Security*, **42**, 165-176.
<https://doi.org/10.1016/j.cose.2013.12.003>
- [22] Bronfenbrenner, U. (1979) *The Ecology of Human Development: Experiments by Nature and Design*. Harvard University Press.
- [23] Tversky, A. and Kahneman, D. (1974) Judgment under Uncertainty: Heuristics and Biases. *Science*, **185**, 1124-1131. <https://doi.org/10.1126/science.185.4157.1124>
- [24] Muraven, M. and Baumeister, R.F. (2000) Self-Regulation and Depletion of Limited Resources: Does Self-Control Resemble a Muscle? *Psychological Bulletin*, **126**, 247-259. <https://doi.org/10.1037/0033-2909.126.2.247>
- [25] Wash, R. (2010) Folk Models of Home Computer Security. *Proceedings of the Sixth Symposium on Usable Privacy and Security*, Redmond, 14-16 July 2010, 1-16.
<https://doi.org/10.1145/1837110.1837125>
- [26] Eppler, M.J. and Mengis, J. (2004) The Concept of Information Overload: A Review of Literature from Organization Science, Accounting, Marketing, MIS, and Related Disciplines. *The Information Society*, **20**, 325-344.
<https://doi.org/10.1080/01972240490507974>
- [27] Florencio, D. and Herley, C. (2007) A Large-Scale Study of Web Password Habits. *Proceedings of the 16th International Conference on World Wide Web*, Banff, 8-12 May 2007, 657-665. <https://doi.org/10.1145/1242572.1242661>
- [28] Pfleeger, S.L. and Pfleeger, C.P. (2012) *Analyzing Computer Security: A Threat/Vulnerability/Countermeasure Approach*. Prentice Hall.
- [29] Starcke, K. and Brand, M. (2012) Decision Making under Stress: A Selective Review. *Neuroscience & Biobehavioral Reviews*, **36**, 1228-1248.
<https://doi.org/10.1016/j.neubiorev.2012.02.003>
- [30] Witte, K. and Allen, M. (2000) A Meta-Analysis of Fear Appeals: Implications for Effective Public Health Campaigns. *Health Education & Behavior*, **27**, 591-615.
<https://doi.org/10.1177/109019810002700506>
- [31] Rogers, R.W. (1975) A Protection Motivation Theory of Fear Appeals and Attitude Change1. *The Journal of Psychology*, **91**, 93-114.
<https://doi.org/10.1080/00223980.1975.9915803>
- [32] Workman, M. (2007) Gaining Access with Social Engineering: An Empirical Study of the Threat. *Information Systems Security*, **16**, 315-331.
<https://doi.org/10.1080/10658980701788165>
- [33] National Institute of Standards and Technology (2024) *The NIST Cybersecurity Framework (CSF) 2.0*. U.S. Department of Commerce.
<https://doi.org/10.6028/NIST.CSWP.29>
- [34] Bulgurcu, B., Cavusoglu, H. and Benbasat, I. (2010) Information Security Policy Compliance: An Empirical Study of Rationality-Based Beliefs and Information Security Awareness1. *MIS Quarterly*, **34**, 523-548. <https://doi.org/10.2307/25750690>
- [35] Vakhula, O., Kurii, Y., Opirskyy, I. and Susukailo, V. (2024) Security as Code Concept for Fulfilling ISO/IEC 27001:2022 Requirements. *CEUR Workshop Proceedings*,

- Vol-3654: Proceedings of the Workshop Cybersecurity Providing in Information and Telecommunication Systems (CPITS 2024)*, Kyiv, 28 February 2024, 59-72.
<https://doi.org/10.55056/ceur-ws.org/vol-3654/paper6.pdf>
- [36] Shaw, E., Ruby, K.G. and Post, J.M. (1998) The Insider Threat to Information Systems: The Psychology of the Dangerous Insider. *Security Awareness Bulletin*, No. 2, 1-10.
 - [37] Cappelli, D.M., Moore, A.P. and Trzeciak, R.F. (2012) *The CERT Guide to Insider Threats: How to Prevent, Detect, and Respond to Information Technology Crimes*. Addison-Wesley.
 - [38] Morewedge, C.K., Yoon, H., Scopelliti, I., Symborski, C.W., Korris, J.H. and Kassam, K.S. (2015) Debiasing Decisions: Improved Decision Making with a Single Training Intervention. *Policy Insights from the Behavioral and Brain Sciences*, **2**, 129-140.
<https://doi.org/10.1177/2372732215600886>
 - [39] Ericsson, K.A., Krampe, R.T. and Tesch-Römer, C. (1993) The Role of Deliberate Practice in the Acquisition of Expert Performance. *Psychological Review*, **100**, 363-406. <https://doi.org/10.1037/0033-295x.100.3.363>
 - [40] Hofstede, G. (2001) *Culture's Consequences: Comparing Values, Behaviors, Institutions, and Organizations across Nations*. 2nd Edition, Sage.
 - [41] Blythe, J.M. and Johnson, S.D. (2018) The Consumer Security Index for IoT: A Protocol for Developing an Index to Improve Consumer Decision Making and to Incentivize Greater Security Provision in IoT Devices. *Living in the Internet of Things: Cybersecurity of the IoT—2018*, London, 28-29 March 2018, 1-8.
<https://doi.org/10.1049/cp.2018.0004>
 - [42] Talent, U.A.S.C. (2023) National Cyber Workforce and Education Strategy.
<https://covacci.org/wp-content/uploads/2023/08/National-Cyber-Workforce-and-Education-Strategy-NCWES-2023.07.31.pdf>