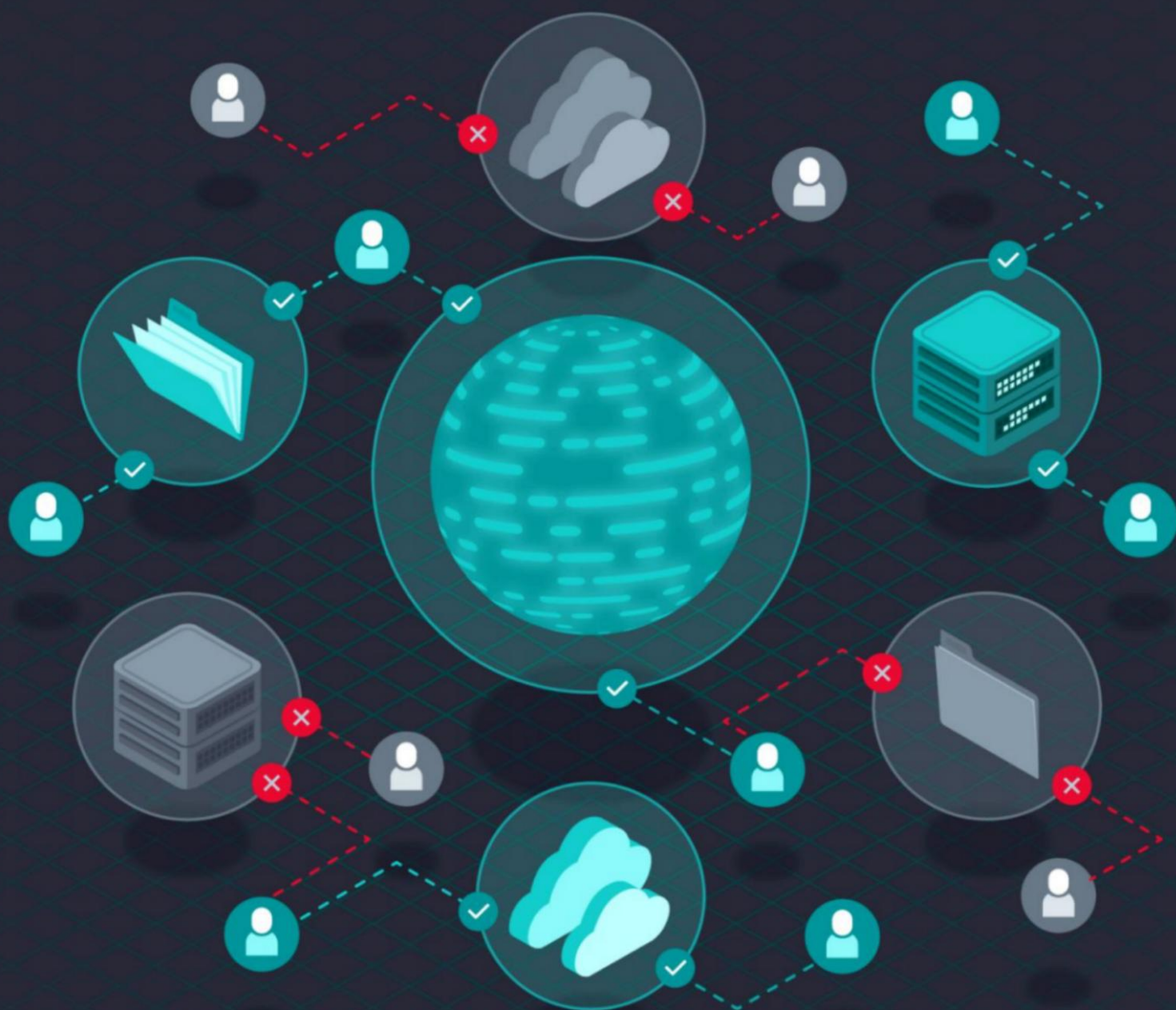


零信任指导原则



零信任工作组（Zero Trust Working Group）的永久官方地址为

<https://cloudsecurityalliance.org/research/working-groups/zero-trust>

@2025 云安全联盟大中华区—保留所有权利。你可以在你的电脑上下载、储存、展示、查看及打印，或者访问云安全联盟大中华区官网（<https://www.c-csa.cn>）。须遵守以下：（a）本文只可作个人、信息获取、非商业用途；（b）本文内容不得篡改；（c）本文不得转发；（d）该商标、版权或其他声明不得删除。在遵循中华人民共和国著作权法相关条款情况下合理使用本文内容，使用时请注明引用于云安全联盟大中华区。



云安全联盟

创立于2009年,作为世界领先的独立、权威国际产业组织,致力于定义和提高业界对云计算和下一代数字技术安全最佳实践的认识和全面发展。



云安全联盟大中华区

在香港注册并在上海登记备案的国际NGO组织,旨在立足中国,连接全球,推动中国数字安全技术标准与产业的发展及国际合作。



4 大区

大中华区、美洲区、
欧非区、亚太区



180+ 分会

英国、法国、加拿大、旧金山、马来西亚等覆盖50多个国家和地区



2.5K+ 成员单位

世界500强科技公司、安全厂商、中小企业、研究机构



20W+ 社区专业人员

研究工作组专家、社区志愿者、从业人员、CSA认证学员



前沿研究

#云安全 #AI安全
#零信任 #数据安全
#5G安全 #区块链安全
#量子安全 #物联网安全
#金融安全 #医疗安全
#智能座舱安全
#关键基础设施安全
.....



培训与认证

CCSK 云安全知识认证
CDSP 数据安全认证专家
CAISP 人工智能认证专家
CZTP 零信任认证专家
CCPTP 云渗透测试认证专家
CCAK 云计算审计知识认证
.....



会议活动

CSA summit@RSAC
CSA GCR Congress
CSA研讨会
AI For GOOD峰会
.....



评估与认证

AI STR AI安全、可信、负责任认证
STAR 云安全评估认证
CAST 云应用安全可信认证
CNST 云原生安全可信认证
.....

1000+研究成果

10W+认证学员

1000W+传播量

成员单位(部分)



企业合作微信号:csagcr



认证培训微信号:CSAlynn



邮箱:info@c-csa.cn



致谢

《零信任指导原则》由 CSA 云安全联盟专家编写，并由 CSA 大中华区零信任工作组完成翻译并审校。（以下排名不分先后）：

中文版翻译专家组

翻译组成员：

杨正权 汪海 吴满 赵晨曦

审校组成员：

陈本峰 陈珊 卜宋博

研究协调员：

郑元杰

英文版本编写专家

主要作者：

Alex Sharpe

贡献者：

Robin Basham

Madhav Chablani

Frank DePaola

Jonathan Flack

Sai Honig

Shamik Kacker

Andrea Knoblauch

Alice Muravin

Rajesh Murthy

Denis Nwanshi

Lars Ruddigkeit

Paul Simmonds

Nelson Spessard

Bernd Wegmann

Heverin Joy Williams

Lauren Wise

审阅者:

Sam Aiello

Jason Garbis

Brett James

Yves Le Gelard

Jennifer Minella

Chandrasekaran Rajagopalan

Aaron Robel

Michael Roza

Vaibhav Malik

Meghana Parwate

Sven Olensky

Annabelle Lee

Himanshu Sharma

CSA 职员:

Erik Johnson

Stephen Lumpe

Stephen Smith

Marina Bregkou

目 录

致谢	4
摘要	8
执行摘要	9
简介	11
目标受众	11
指导原则	12
1. 以终为始（业务/任务目标）	12
2. 避免过度复杂化	13
3. 产品并非优先事项	15
4. 访问是一个有意的行为	16
5. 从内到外，而不是从外到内	17
6. 安全漏洞不可避免	19
7. 了解您的风险承受能力	21
8. 确保高层的支持和倡导	25
9. 灌输零信任文化	26
10. 开始小规模并专注于快速赢得成果	27
11. 持续监控	28
12. 实用参考资料	29
13. 推荐阅读	29

摘要

零信任（ZT）是一种战略思维模式，对于组织而言，在数字化转型以及其他旨在增强组织安全性与弹性的工作中，采用零信任理念极具价值。然而，由于安全行业内信息混杂且缺乏统一标准，零信任常很容易被误解也容易被过度复杂化。事实上，零信任是基于一些长期存在的原则，随着我们工作和生活方式的改变，这些原则变得愈发关键，例如远程办公、对第三方依赖的增加、云计算的采用，以及机器学习（ML）、自然语言处理（NLP）和大语言模型（LLM）等人工智能（AI）技术的广泛和加速应用。本文档旨在填补这些空白，通过梳理基本原则，包括诸如最小权限原则、职责分离原则和分段原则等既定的信息安全（InfoSec）原则，来阐明零信任理念。这些指导原则将在所有零信任支柱、不同的应用场景、环境和产品中都将保持一致。随着行业的发展，本指南也会不断演进。

执行摘要

零信任是一种简单的信息安全（InfoSec）方法，但常常被误解和过度复杂化。如果能够正确理解，零信任的理念和策略将成为组织提升安全性、增强弹性并指导数字化转型的有力工具。这些工具也可以与 AI 结合，实现双赢：AI 可以帮助实施零信任计划，而零信任原则也应用于保护 AI 的使用，包括 AI 模型和训练数据。本文档旨在清晰阐述零信任的概念，并提供在规划、实施和运营零信任时需牢记的指导原则。

在传统模式下，信息安全主要依赖于技术控制，其安全模型基于将资产集中并置于受控的物理边界内进行防护。然而，这种情况已不复存在。零信任认识到人、流程、组织和技术之间的整体关系，并指出仅依靠技术控制已不足够。过去，用户仅因其位于组织边界内就被认为是“可信的”。零信任颠覆了这一概念，要求无论用户位于何处，在授予资产访问权限之前都必须进行验证。

零信任利用长期以来一直存在的原则，例如“永不信任，始终验证”、最小权限原则以及网络分段实践等，以提高网络安全水平，降低总拥有成本（TCO），减少事件造成的损害，并加快恢复速度。通过将零信任原则与现有安全实践相结合，组织可以为在复杂的分布式环境中保护资产奠定坚实的基础。这种主动的方法提升了安全态势，并减少了因威胁环境不断演变而带来的潜在风险。

零信任也认识到安全漏洞不可避免。为了增强弹性，零信任提供了一种控制“爆炸半径”的方法，以减少入侵的影响，同时促进快速恢复。这些技术还增加了攻击者的工作量和成本，从而进一步降低了安全事件发生的可能性。

近年来，零信任的兴起得益于新的商业模式、云计算和 AI 的广泛应用，以及新的政府要求。美国已通过行政命令要求所有联邦机构实施零信任^[1]。在全球范围内，《数字操作弹性法》（DORA）^[2]和欧盟的《网络和信息安全指令》（NIS2）^[3]等举措也推动了零信任的采用。零信任通过一系列适用于所有零信任计划的基本原则，提供了必要的安全保障。本文档概述了这些基本原则，旨在为任何零信任计划提供指导。

-
1. <https://www.whitehouse.gov/briefing-room/presidential-actions/2021/05/12/executive-order-on-improving-the-nations-cybersecurity/>
 2. <https://www.digital-operational-resilience-act.com/>
 3. [https://www.europarl.europa.eu/thinktank/en/document/EPRS_BRI\(2021\)689333](https://www.europarl.europa.eu/thinktank/en/document/EPRS_BRI(2021)689333)

简介

各组织广泛采用零信任理念来推动数据和网络安全管理实践的转型。许多零信任管理概念已经出现，包括原则、理念、支柱、架构计划和框架。然而，这种演变是一个长期的过程，无法通过单一项目（业务、运营、技术）或特定产品来实现零信任转型。零信任是一种成熟的方法论，旨在增强分布式架构中关键资产的保护能力。零信任实施需要提前规划，且所有关键利益相关者必须理解每个零信任实践过程都是独特的。与业务目标的对齐程度越高，零信任实践成功的可能性就越大。

许多组织已经调整其运营模式，以支持云化部署、远程办公和 AI 等场景。传统的安全实践已无法充分应对由此产生的新风险格局。希望提升网络弹性的组织不能再仅仅依赖外部边界保护层或单纯的技术控制来缓解网络风险。网络威胁形势不断演变，已超出传统堡垒模式的防御能力。

同时，需要保护的范围也扩大了。我们不再仅仅关注 IT 资产和数据，而是将保护范围扩展到了设备、工作负载、应用程序和 IT 之外的业务流程。这通常被称为数据、应用程序、资产和服务（简称 DAAS）。

随着 AI 的日益普及，零信任原则也被用于保护 AI 模型。例如，限制谁可以访问、修改或训练模型，这些都是零信任原则的典型应用。

通过将安全架构与业务运营模式对齐，组织可以在不影响业务流程的前提下，在提供适当安全保障的同时，推动业务转型。当零信任被视为基础理念时，它还能够支持组织的其他工作，例如隐私保护、合规性和风险管理。

本文档提供了组织在规划或启动零信任时可以参考的指导原则。

目标受众

本文档的主要受众是信息保护从业人员及其高管层。作为指导原则，本文档的内容适用于所有零信任计划。我们认为这些指导原则将被行业组织和标准组织采用，以构建零信任知识体系（BOK）。

指导原则

零信任并不是一个独立的概念或技术。相反，它是一种全面的安全战略和方法，涵盖各种原则、策略和技术，旨在应对不断变化的威胁形势以及传统基于边界的安全模型的局限性。需要特别指出的是，指导原则是稳定不变的。然而，这些原则对每个组织的意义和价值会因国家、行业和具体组织的不同而有所差异。

以下指导原则旨在帮助从业者保持方向并管理零信任建设实践。

（正文内容如下）

1. 以终为始（业务/任务目标）

零信任是一种与传统边界安全模型（堡垒模型）截然不同的安全范式。在堡垒模型中，内部人员被认为是可信的，而外部则被视为风险来源。这种模型在企业集中在一个静态地点时运作良好。而如今，组织存在于分布式的，通常是全球化的生态系统中。零信任旨在将安全架构与组织的分布式工作模式和技术模式对齐，在这种模式下，不再有“内部”和“外部”之分。供应链风险管理（SCRM）和第三方风险管理（TPRM）是分布式生态系统中的关键领域，零信任原则在这些领域中发挥着基础性作用，是零信任建设的典型应用场景。

以终为始意味着要有清晰的愿景和目标，从而更快地实现成果，同时避免资源浪费。这种结果导向的思维方式也促进了业务战略、运营模式和安全架构之间的协调与对齐。

常见的预期结果包括：

- 在合理的风险下创造价值，无论是通过推出新产品、进入之前无法触及的市场，还是获得未知的竞争优势。
- 通过为所有合规、网络弹性和隐私需求奠定基础，降低合规成本。

- 降低安全事件的影响（例如成本）。
- 减少 IT 的复杂性，减轻流程负担。
- 降低总体拥有成本（TCO）。
- 为第三方风险管理（TPRM）奠定基础。
- 建立更具弹性的治理、风险管理和合规（GRC）计划，以更好地防御当前和未来的威胁。



2. 避免过度复杂化

人们很容易忘记，零信任的核心是一系列长期存在的原则，这些原则通过将安全架构与我们的工作和生活方式对齐来提升安全性。向零信任解决方案演进并不像看起来那么复杂，尤其是当您将其作为一系列有计划的步骤逐步推进，并优先保护最关键资产时。一旦最关键资产得到保护，您可以根据资产的重要性依次处理其余资产。

已实施和测试的安全控制措施，包括预防性、检测性和纠正性（或反应性）控制措施，构成了零信任的基础。这些基本原则对零信任的成功至关重要：

- 最小权限访问控制概念（例如，预防性）
- 职责分离（例如，预防性）
- 分段/微隔离（例如，预防性）
- 日志记录和监控（例如，检测性）
- 配置漂移修复（例如，纠正性/反应性）



零信任基于这些核心原则，帮助组织从传统的堡垒模型转向现代组织中常见的分布式模型。它还为控制措施的精细化和敏感性提供了机会，例如通过增加持续的身份验证和授权（如预防性控制）、用户与实体行为分析（UEBA）（如检测性控制），以及动态策略执行点（如纠正性/反应性控制）。

人工智能（AI）带来的生产力提升，有望促进监控、捕捉配置错误，并检测权限过高的问题。更先进的 AI 分析技术，有望实现更快的检测、根

本原因分析以及可疑模式的识别。

控制类型不需要为零信任重新设计或过度复杂化。相反，现有控制措施的速度、性能和灵活性可以得到优化。如今，人工智能（AI）被用于提供优化，能够帮助降低复杂性、减少总拥有成本（TCO）并缩短响应时间。

安全基础构成了零信任的核心。例如，最小权限原则和职责分离就是典型的例子。确保所有用户（员工、合同工和供应商）都得到唯一标识，并定期审查他们的权限，根据需要更新，也是一个关键点。身份生命周期管理至关重要。迄今为止，这些基础概念的实施和维护主要是手动的，自动化程度较低。人工智能（AI）的使用有望提供至少一定程度的自动化和检测能力。

引入分段和微隔离有助于加强上述控制措施，并减少事件的影响。自动化可以简化常规任务中所需的手动流程量。

3. 产品并非优先事项

从历史上看，所有与网络相关的事务都属于技术专家的领域。用技术解决技术问题是自然的选择，通常这会导致产品购买，并常常伴随着咨询服务。在许多方面，零信任更多关注的是人、流程和组织维度，而不是技术本身。

过度依赖产品而不考虑人、流程和组织的战略不会成功。仅依赖产品购买来实现零信任之旅并不是真正的零信任之旅。如果您首先解决其他维度的问题，您将更好地理解需求，从而支持更强大的长期零信任策略，并增加选择合适产品（如有必要）的可能性。

人工智能（AI）带来的生产力提升，以及责任向“安全设计与默认”^[4]等举措的转移，有望进一步减少对额外产品和资源的依赖。

4. <https://www.cisa.gov/resources-tools/resources/secure-by-design>

4. 访问是一个恶意的行为

零信任的一个关键区别在于它不依赖于物理或网络边界。在传统模型中，例如堡垒模型，用户被授予对网络的访问权限通常意味着他们可以访问其他资产。这种理念的产生是由于当时技术的局限性，并且当时我们生活在一个可以将资产集中在边界内的世界中。

零信任将责任转移给产品供应商和服务提供商，正如“安全设计与默认”^{5]}以及《国家网络安全战略》^{6]}中所体现的。这种转变减轻了组织在实施补偿性控制和其他安全实践（如漏洞管理）方面的负担，并允许资源重新分配到更具战略性和主动性的活动中。

如今，组织处于一个全球化的环境中，远程员工增多，对云等第三方的依赖增加，供应链也变得越来越复杂。在许多方面，组织越来越依赖外部资源。

技术已经发展到我们可以超越物理或网络边界的程度。零信任利用了这些技术进步，使我们能够更好地识别用户，并更频繁地做出更细粒度的访问控制决策。

在当今世界，身份必须经过明确验证，只有在授权通过验证过程后，才会授予访问权限。

历史上，我们主要依赖 IT 组织来决定谁可以访问哪些资产。近年来，越来越多地依赖业务所有者（例如数据所有者）和流程所有者来决定谁可以访问什么、访问多久。IT 组织越来越被视为保管者（例如数据保管者）。此外，随着技术的进步，以及我们能够收集用户和实体行为分析（UEBA）、设备、网络 and 应用程序信号，组织正在寻求利用人工智能（AI），不仅学习在特定上下文中安全的访问权限，还能实时响应。

5. <https://www.cisa.gov/resources-tools/resources/secure-by-design>

6. <https://www.whitehouse.gov/wp-content/uploads/2023/03/National-Cybersecurity-Strategy-2023.pdf>

5. 由内到外，而不是由外到内

采用从内到外的策略，您的企业政策将从“我们要防御什么？”转变为“我们要保护什么？”。

传统的安全模型依赖于强大的外部边界。这些模型假设内部的任何人都可信的，而外部的任何人都是不可信的。随着过去二十多年的“去边界化”^[7]发展，更多的人员和资产位于边界之外。这意味着传统的“外防内信”安全理念已不再适合我们当前的工作和生活方式。由于没有任何组织拥有无限的资源（如时间、资金、精力），我们需要将精力集中在最有价值的地方。

资产的价值是指导我们优先级排序的依据。如果您有业务影响评估（BIA），请从这里开始。如果没有，请进行资产盘点并按价值对资产进行分类。使用从高到低的资产价值排序来指导您的工作。

一旦明确了需要保护的资产及其关系，您就可以识别出保护面和攻击面。

在现代组织中，我们试图保护的通常是数据、应用程序、资产和服务。它们统称为 DAAS。

7. https://en.wikipedia.org/wiki/Jericho_Forum

数据	这是一类敏感数据，如果被外泄或误用，可能会使组织陷入麻烦。通常，数据之所以被认为是敏感的，可能是因为某个第三方（如监管机构）声明（认定）它是敏感的，或者因为它是你的 知识产权或组织运营所必需的数据。
应用程序	软件、硬件和通常为满足一组需求而协作的基础设施的集合。零信任原则同样适用于 AI、模型、数据、输入和输出的保护和操作。
资产	提供组织拥有或控制的价值资源。资产可能包括信息技术（IT）、操作技术（OT）或物联网（IoT）设备，包括销售点（POS）终端、SCADA 控制、制造系统和医疗设备。
服务	通过提供业务和技术专业知识来实现组织期望结果的应用，而无需拥有具体的成本和风险。在现代企业中，服务可以是基于云的，如软件即服务（SaaS），也可以是应用程序之间的常见用例，例如应用程序编程接口（API）或域名系统（DNS）。

Kipling 方法是一种用于制定政策的标准工具。它对于确保安全战略与业务模型的一致性至关重要。

这些问题只能通过资产所有者（业务部门）与资产保管者（通常是 IT 部门）之间的合作来解答。

- 谁可以访问资产？他们被允许做什么？我能确保他们确实是他们声称的身份吗？
- 他们试图访问什么资产？他们被允许采取什么行动？
- 被允许的访问何时开始？何时结束？是否仅在某些特定时间段内允许访问？
- 资产位于何处？它是否只能从某些特定位置访问？是否有某些位置不允许访问该资产？
- 为什么这个用户需要访问这个资产？保护资产的原因是其敏感性。敏感性是否由合规要求定义？
- 是否存在访问资产的限制方式？

假设组织已经收集了与本节早前详细描述的数据，人工智能（AI）可以用来生成这些政策的初步版本，但仍需经过政策作者的进一步审查和审定。

6. 安全漏洞不可避免

假设你可以完全保护自己免受网络风险是不现实的。防御者无法堵住所有漏洞，而攻击者只需找到一个漏洞即可。这有时被称为“防御者的困境”。零信任实施的核心是在授予资产访问权限之前直接验证声明的身份和许可的访问权限。

传统模型主要依赖于物理和技术控制。组织也依赖于他们的能力来阻止恶意行为者。随着世界变得更加数字化，传统的防御墙变得更加脆弱，世界认识到漏洞是不可避免的。它们通常由伪装成有效内部人员的外部攻击者促成，有时也由实际的内部人员（例如内部威胁）促成。零信任的作用是减少漏洞发生的可能性，降低其影响，并促进更快的恢复。我们通过实施更健壮的访问控制、更注重检测潜在事件，并规划事件响应和恢复来实现这一点。

一旦理解到大多数事件根本上是人为问题，就会理解漏洞是不可避免的。与其专注于保持安全，实践者的心态转向了弹性。减少事件影响范围的关键是分段和微分段。每个分段通过限制在组织内部的横向移动能力和限制恶意行为的传播来减少可能性和影响。例如，如果恶意行为者接管了有效用户的账户，分段意味着他们无法访问有效用户范围之外的区域。在恶意软件（例如勒索软件）的情况下，单台机器的感染不会在组织内部传播。

在更传统的安全模型中，一旦你进入了边界内，你就被信任，并且可以在这个边界内自由移动。旧的安全架构就像一个“硬壳软核”的吉百利巧克力蛋——外部防护严密，但是内部缺乏纵深防御。想象一下，间谍在进入社区后可以肆无忌惮地活动。

在零信任模型中，用户和设备在任何地方都是不被信任的。在获得任何资产访问权限之前，必须进行审查。想象一下，一个旧欧洲城市，街道曲折，房子挨着房子。邻居彼此认识。如果你走到这些街道上，每个人都会怀疑并进行审查。你来到这里做什么？你的目的是什么？你是小偷吗？游客吗？还是来这里做生意的？这种审查不是偶尔发生，而是你去的每一个地方都会发生。零信任也是如此。每次请求访问资产时，都会验证身份、验证授权访问，并且交互被记录。审查的程度与资产的价值和风险环境相适应。

人工智能工具有望进一步实现数据审查和分析的自动化。

这种心态的转变带来了零信任旅程中的关键成果。首先，它限制了漏洞发生时的影响范围（并且它们会发生）。其次，它减少了黑客在组织内横向移动的能力。第三，它通过限制单一事件可以破坏的资产来减少影响。

从董事会和高级领导层的角度来看，预测性可能是关键成果。当用户发生漏洞时，你知道潜在影响是有限的——不再是整个组织。你可以快速识别风险所在。你的可能性不再是 100%，影响也不再是不可计算的。

有效地关联身份、行为和资产的能力极大地增强了数据防泄漏（DLP）的工作。没有明确的边界，DLP 不再像以前监控进出流量那样简单。通过利用零信任原则和创新的数据发现机制，我们可以在潜在漏洞发生之前检测它们，在漏洞发生时更快地捕捉它们，并减少管理和操作数据泄漏任务的复杂性。

人工智能的潜在优势将有所帮助，但无法完全解决问题。

关键是要持续验证，并以与资产价值相称的方式进行验证。寻找用户行为中的异常，并假设恶意行为者想要你的资产，这是其中重要组成部分。如今，AI 和 UEBA 相结合可以用于自动化身份、行为和资产数据的关联过程，以检测和响应异常活动，并做出上下文相关的数据防泄漏决策。

7. 了解您的风险承受能力

风险承受能力是一个被广泛接受的风险管理概念。风险承受能力是组织在追求其目标时愿意接受的风险水平。固有风险是指在采取行动（例如，风险处理）之前存在的风险水平。目标是将风险处理到低于风险承受能力的水平。这称为可接受风险。

人工智能的广泛使用和加速采用为理解组织的风险承受能力提供了新的维度。人工智能既带来了新的风险，同时也放大了其他风险的影响并增加了其发生的可能性。值得注意的是，人工智能更加强调 CIA（Confidentiality - 保密性、Integrity - 完整性、Availability - 可用性）三要素中的完整性部分。在模型的完整性、训练数据、输入、输出和模型访问等方面应给予谨慎考虑。

零信任的指导原则当然适用于 CIA 三要素中的完整性部分以及人工智能。所使用的工具和技术可能需要进行调整，但原则保持不变。例如：

- **人工评估：**这是验证输出的最可靠方法之一。人工评估者可以评估生成内容的质量、相关性和适当性。然而，这种方法可能耗时且主观。
- **自动化指标：**有几种自动化指标用于评估生成文本的质量，例如用于机器翻译的 BLEU^[8]、ROUGE^[9] 和 METEOR^[10]，以及用于语言模型的困惑度（Perplexity）。这些指标将生成的输出与参考输出进行比较。然而，它们可能并不总是与人类判断一致，尤其是在创造性任务中。
- **A/B 测试：**在实际系统中，可以使用真实用户测试模型的不同版本。用户参与度或满意度更高的版本被认为是更优的。
- **安全性和公平性检查：**确保人工智能不会生成有害、偏见或不适当的内容非常重要。这可以通过在精心策划的数据集上进行预训练和微调，并结合基于人类反馈的强化学习来实现。

8. <https://en.wikipedia.org/wiki/BLEU>

9. [https://en.wikipedia.org/wiki/ROUGE_\(metric\)](https://en.wikipedia.org/wiki/ROUGE_(metric))

10. <https://aclanthology.org/W05-0909.pdf>

- 一致性检查：对于某些任务，可以检查生成输出的一致性。例如，在问答任务中，可以用不同的方式询问相同的问题，并检查答案的一致性。

组织并非单一的整体。组织的不同部分通常具有不同的可接受风险水平，这些通常被称为风险容忍度。例如，金融服务公司的风险投资部门会比同一机构的固定收益部门容忍更高的风险。虽然这种区别很重要，但零信任的指导原则同样适用于风险偏好和风险容忍度。为了可读性，讨论中不做区分。

零信任通过实施控制措施将固有风险降低到可接受水平，这些控制措施旨在减少风险发生的可能性^[11]或降低其影响^[12]^[13]，或两者兼而有之。

CIA 三要素是一个被广泛接受的信息安全（InfoSec）模型，旨在帮助组织评估资产潜在的损害（例如，影响）。为了完整起见，这里同时包含 NIST 和 ISO 的定义。



1. 保密性（Confidentiality）：保持对信息访问和披露的授权限制，包括保护个人隐私和专有信息的手段。^[14]即信息不会被未经授权的个人、实体或过程获取或披露。^[15]数据泄露是资产保密性丧失的一个例子。

11. <https://csrc.nist.gov/glossary/term/likelihood>

12. <https://csrc.nist.gov/glossary/term/impact>

13. ISO 31000, Risk Management, uses the term “consequences”

14. <https://csrc.nist.gov/glossary/term/confidentiality>

15. <https://www.iso.org/standard/14256.html>

2.完整性（Integrity）^{【16】}：防止信息的非法修改或破坏，包括确保信息的不可抵赖性和真实性。^{【17】}即数据没有以未经授权的方式被修改或销毁。^{【18】}完整性的丧失是三者中最难以把握的。一般而言，如果它导致你的信任丧失，那么就是对完整性的攻击。例如，文件被计算机病毒感染。深度伪造（Deep Fakes）、AI 幻觉和各种形式的欺诈都涉及到完整性的丧失。

3.可用性（Availability）： 确保信息的及时和可靠的访问和使用。^{【19】}

可根据授权实体的要求访问和使用的属性。^{【20】}

勒索软件（Ransomware）可能是资产可用性丧失最为显著的结果之一。勒索软件对数据进行加密，使业务无法访问这些数据。当像水这样的公共设施访问被拒绝时，这也是对可用性的攻击。

所有组织都需要确定其风险偏好。这一决定远远超出了零信任的范围。作为这一过程的一部分，了解已经达成的协议、零信任的角色以及组织常用的工具是明智的做法。最常见的工具可能是风险登记册（Risk Register）。网络风险是少数几种业务必须管理的风险之一，它可能影响到许多其他领域的风险，使得完全量化变得困难。许多组织选择使用定性措施来评估风险，例如“非常高”、“高”、“中”、“低”。

16. ISO only has a definition for Data Integrity not Integrity

17. <https://csrc.nist.gov/glossary/term/integrity>

18. <https://www.iso.org/standard/14256.html>

19. <https://csrc.nist.gov/glossary/term/availability>

20. <https://www.iso.org/standard/14256.html>

其他组织则被要求或选择遵循某些标准，例如欧盟的数字运营弹性法案（DORA）^[21]、网络和信息安全指令（NIS2）^[22]、NIST SP 800-53^[23]，或联邦金融机构检查委员会（FFIEC）制定的网络安全评估工具（FFIEC 的 CAT）^[24]。

一些组织选择采用不同的量化技术来计算风险价值（VaR）^[25]。其中最著名的之一是信息风险的因子分析（FAIR）^[26]。无论您采取哪种策略，原则保持不变：零信任的核心目的是将风险降低到可接受的水平。

组织不可避免地必须不断演变，否则将面临消失的风险。威胁环境也在不断变化，迫使组织不断重新评估其风险偏好。这些变化可能会揭示新的漏洞或新的利用方式。更复杂的是，总会存在一些组织无法充分预测或准备的未知因素。勒索软件就是一个典型的例子。尽管这一概念在社区中早已存在，但由于在加密货币普及之前不够实际，因此未能成为主流。总的来说，这些类型的风险被称为“未知的未知（unknown-unknowns）”。

零信任的基本阻止和应对策略在很大程度上有助于使组织在面对这些演变和未知的未知时具有未来保障。随着新的威胁出现、新的漏洞的产生或之前未知的未知问题浮现时，零信任的核心原则将有助于降低风险的发生可能性或影响，甚至可能同时降低两者。

新冠疫情、地缘政治发展、供应链冲击、高影响的气候变化事件以及对关键基础设施日益增长的网络威胁，已经导致政府和组织在运营弹性方面的范式转变。更加深入地理解和量化风险容忍水平将变得越来越重要。我们已经看到全球其他地区定义了新的立法，例如欧盟的《数字运营弹性法案（DORA）》。这将对组织的风险管理（ERM）实践产生长期影响，组织必须不断创新和适应新的网络威胁。

21. <https://www.digital-operational-resilience-act.com/>

22. [https://www.europarl.europa.eu/thinktank/en/document/EPRS_BRI\(2021\)689333](https://www.europarl.europa.eu/thinktank/en/document/EPRS_BRI(2021)689333)

23. <https://csrc.nist.gov/publications/detail/sp/800-53/rev-5/final>

24. https://www.ffiec.gov/pdf/cybersecurity/ffiec_cat_may_2017.pdf

25. <https://www.investopedia.com/terms/v/var.asp>

26. <https://www.fairinstitute.org/>

这正是零信任发挥关键作用的地方。至少，它将使组织更具韧性并提高灵活性。

8. 获取高层的支持和理解

零信任是一项组织性工作，而不仅仅是技术问题，因此需要组织内部各级的合作才能成功实现。这只能通过适当的高层支持和来自高层的明确传达来实现。必须任命合适的负责人，并保持其信息更新。在理想的情况下，高层领导将积极沟通零信任的重要性，包括与业务战略的一致性、适当的资本分配和企业政策的对齐。

在最理想的情况下，零信任应该由董事会支持。至少，它需要得到高级领导的支持。如果零信任努力仅限于某个业务单元或地理区域，那么最好由业务单元领导或该地理区域内的组织高层（如国家领导）进行支持。

沟通是成功实施零信任的关键。一份结构化的沟通计划可以确保参与者和利益相关者之间的信息流动一致且可追溯，从而有助于指导零信任的落地。

绘制映射利益相关者的图表可以清晰地展示各个角色的职责和责任。通常称为 RACI 图（责任、问责、咨询、通知），它描述了不同角色在完成项目或业务流程中的任务或交付物时的参与情况。领导者应通过全力支持零信任模型并强调其对组织的重要性来树立榜样。他们应定期沟通其重要性。此外，建立一种企业文化至关重要，这种文化认识到网络风险是每个人的责任，而不仅仅是少数人的责任。

9. 灌溉零信任文化

“文化是人们在无人注视时的行为。”^{【27】}

零信任是每个人的责任，而不仅仅是 IT 部门或首席信息安全官（CISO）的职责。将零信任原则融入企业文化中，能够确保组织中每个人每天做出的成千上万的决策都与战略方向一致，同时也有助于抵御日益增多的社会工程攻击。这一点尤为重要，因为 68% 的数据泄露事件与针对人类因素的攻击有关。^{【28】}

将零信任指导原则融入安全以及意识培训是一个很好的起点。恶意攻击者越来越多地利用 AI 技术量身定制攻击。类似地，针对最受攻击的角色量身定制的培训有助于对抗这些攻击，并提高组织的网络安全水平。拥有特权账户访问权限的管理员、高价值目标（如高管）以及处于最高风险的角色（如应付账款）是一个很好的起点。

然而，近期事件突显了培训和意识提升的必要性及其价值，即使这些角色最初看起来并非高风险或高价值。2023 年秋季攻击赌场的勒索病毒攻击就是一个典型例子。恶意攻击者利用了客户服务代表（CSR）乐于助人的自然倾向，欺骗他们通过密码重置提供对高权限账户的访问。虽然我们无法确定通过对 CSR 进行零信任原则的培训是否能避免这些事件，但我们知道，培训和意识提升是组织可以做出的最低成本的网络安全投资。仅阻止一个事件就足以证明这一投资的价值。

什么是零信任文化？零信任文化是一种员工对需要保护的内容及其保护程度有清晰认识的文化。最重要的是，所有员工都理解，访问授权从不暗示，而是一个明确的行为。所有员工和组织业务往来的个人都应知道如何识别可疑活动，并向适当的内部和外部当局报告任何与网络相关的担忧，

27. Herb Keller, the former CEO of Southwest Airlines

28. Verizon 2024 Data Breach Investigations Report

同时理解某些安全控制措施的存在原因。零信任文化是灵活适应的，并不依赖于任何特定的技术或架构。人们被赋予权力，以当前和未来最合理的方式保护资产。

安全曾经是某个特定部门（例如 IT）的职责，但现在已变得无处不在。开发人员可以接受它，业务领导者可以接受它，最终用户可以通过与他们的设备进行无摩擦的互动体验其好处。最重要的是，组织应该理解，零信任使他们能够更智能地应用技术。采用这种文化的风险可能在于，在分散的文化意识还未完全建立之前，过早削弱集中安全部门的角色。

在组织内部灌输零信任文化意味着在组织的所有层级推广对零信任安全模型的理解和接受。

10. 从小规模开始，专注于快速取得成效

由于零信任是一种策略而非特定的产品集，因此，基于其基本概念（例如从内向外设计）开展工作，可以使团队逐步取得成功，而无需大额的前期支出。应识别并优先保护由 DAAS 元素构成的保护面，并根据其规模 and 影响进行排序。

选择一个小型、低成本的保护面作为试点时，更容易获得并维持领导层的支持，因为可以利用其指标来突显安全范式的转变并展示业务价值。

^[29] 向全面实施零信任迈进将带来许多积极的业务成果。重要的是要持续强调对整个组织的好处以及通过所有网络相关变更降低的风险。

试图承担过多任务并追求耗时较长的胜利，可能会使项目陷入困境，并将组织的零信任努力与失败而非成功联系起来。

29. A detailed description of attack surface and protect surface can be found in the Zero Trust Advancement Center Resource Hub hosted by theCloud Security Alliance, <https://cloudsecurityalliance.org/zt/resources/>

11. 持续监控

零信任假设没有任何参与者（例如用户、设备、服务或应用程序）可以被隐式信任。相反，接受声明的身份或授予请求的访问权限是一个明确的行为。所有对组织资源的访问请求都必须经过身份验证和授权验证，就像它们来自未知来源一样，然后才能授予访问权限。即便如此，访问权限也是有限期的（而非永久的）。

鉴于攻击者经常利用合法用户的账户，而恶意内部人员也常常试图超越权限以满足其需求，因此监控和记录事件至关重要。监控对于早期发现潜在恶意行为至关重要，而日志记录则对于识别入侵指标（IOC）、确定影响和收集证据至关重要。监控和日志记录共同促进了持续改进。重要的是要建立监控机制，最终覆盖整个组织的活动。

随着对验证身份的重视增加，以及在授予资产访问权限之前所需的尽职调查，持续监控的机会也理所当然地增加了。历史上，这种深度的数据并不可用，持续监控的需求也不那么明显。技术的进步，包括人工智能，使得这些数据的自动收集、整合和处理成为可能。虽然人工审查的需求不会很快消失，但技术进步增加了警报能力，同时减轻了资源负担，并促进了收集数据的分析。

监控和维护零信任基础设施涉及定期审计访问权限、持续监控网络行为、维护最新的安全补丁、进行风险评估以及加强用户安全意识。

历史上，监控的广度和深度有限，从监控数据中获取洞察和及时分析的能力也受到限制。在这方面，组织也可以利用人工智能，持续监控用户、应用程序、设备和网络行为，并将其与风险评估、风险决策以及针对用户和系统的安全意识培训相结合。

一个常见的误解是零信任没有边界。在当今互联的世界中，边界不再像过去那样清晰或坚固。然而，我们不能因为无法依赖一种能够确保将攻击者拒之门外并仅允许合法用户进入的策略，而放松我们的尽职义务。恰

恰相反，组织有责任定义、监控和控制内部和外部边界，以保护其资产。

12. 实用参考资料

零信任进阶中心资源库，由云安全联盟（Cloud Security Alliance）主办

访问地址：<https://cloudsecurityalliance.org/zt/resources/>

美国联邦零信任资源中心

访问地址：<https://zerotrust.cyber.gov/>

13. 推荐阅读

美国国家安全电信咨询委员会（NSTAC），关于零信任与可信身份管理的总统报告，2022

年 <https://www.cisa.gov/sites/default/files/publications/NSTAC%20Report%20to%20the%20President%20on%20Zero%20Trust%20and%20Trusted%20Identity%20Management%20%2810-17-22%29.pdf>

零信任成熟度模型第 2 版，美国网络安全和基础设施安全局（CISA），2023 年 4 月 <https://www.cisa.gov/zero-trust-maturity-model>

关于加强美国国家安全的白宫行政命令，白宫，2021 年 5 月 12 日，<https://www.whitehouse.gov/briefing-room/presidential-actions/2021/05/12/executive-order-on-improving-the-nations-cybersecurity/>

关于推进零信任成熟度的 NSA 指导新闻稿，美国国家安全局（NSA），2023 年 <https://www.nsa.gov/Press-Room/Press-Releases-Statements/Press-Release-View/Article/3328152/nsa-releases-recommendations-for-maturing-identity-credential-and-access-manage/>

关于推进零信任成熟度的 NSA 指导，美国国家安全局（NSA），2023 年 3 月 https://media.defense.gov/2023/Mar/14/2003178390/-1/-1/0/CSI_Zero_Trust_User_Pillar_v1.1.PDF

零信任架构，美国国家标准与技术研究院（NIST），特别出版物 800-207，2020 年 <https://csrc.nist.gov/publications/detail/sp/800-207/final>

CSA GCR

Cloud Security Alliance Greater China Region



扫码获取更多报告