

Securing Spaces

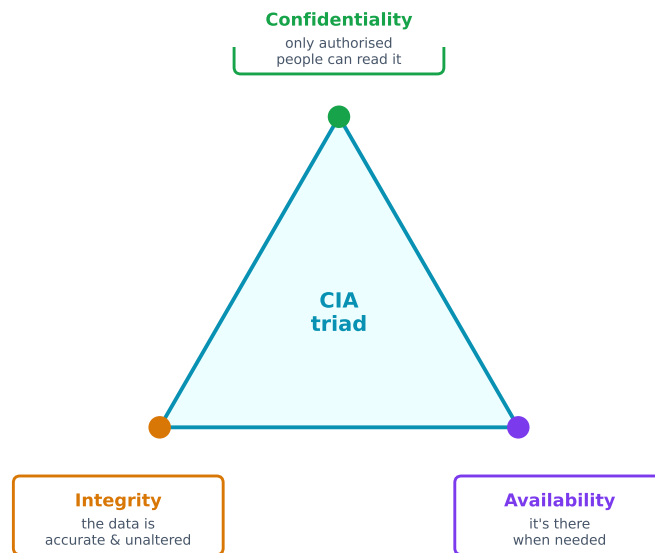
AP Cybersecurity

Cyber Foundations

Before defending a system, you need a shared language. This section builds it.

Every security control protects at least one part of the **CIA triad** 中央情报三要素 - the three goals of security:

- **Confidentiality** 保密性 - only authorised people can read the data.
- **Integrity** 完整性 - the data is accurate and unaltered.
- **Availability** 可用性 - the data and services are there when needed.



The CIA triad: the three goals every security control supports

Attacks come from different adversaries, classified by their goals. A **script kiddie** 脚本小子 reuses tools built by others for greed or recognition; a **hacktivist** 黑客活动分子 acts for a political, social, or personal cause; an **insider** 内部人员 already holds legitimate access and may act from revenge or greed; a **cyberterrorist** 网络恐怖分子 disrupts critical infrastructure like a power grid or water plant; and **transnational criminal organisations** 跨国犯罪组织 chase money through ransomware and stolen data.

Most attacks unfold in **phases** 阶段: **reconnaissance** 侦察 (gathering information, often from public **OSINT** 公开来源情报 sources), initial access, persistence, **lateral movement** 横向移动 (spreading to more systems by escalating privileges), taking action on the goal, and evading detection. Naming the phase an attacker has reached helps a defender choose the right response.

Social engineering: the eight tactics

Most attacks begin not with code but with **social engineering** 社会工程学 - psychological tricks that manipulate a person into doing what the adversary wants. The exam names **eight** tactics, and expects you to identify which one a scenario shows:

Tactic	The trick
Pretexting 借口	inventing a believable reason to make contact ("I'm from IT, verifying your account")
Authority 权威	posing as someone powerful, or relaying "the boss's" instructions
Intimidation 恐吓	threatening negative consequences if a demand is not met
Consensus 从众	claiming everyone else is already doing it, to create social pressure
Scarcity 稀缺	inventing limited availability ("only 2 left")
Familiarity 熟悉	pretending to be, or to know, someone close to the target
Urgency 紧迫感	imposing a tight deadline so the target acts before thinking

The common thread is that all eight bypass a target's judgement by triggering an automatic emotional response - fear, trust, haste, or the wish to fit in. The defence is the same each time: **verify through a separate, trusted channel** before acting.

A **risk** 风险 appears when a **threat** 威胁 can exploit a **vulnerability** 漏洞 to compromise an **asset** 资产 (anything valuable - data, money, hardware, reputation). We **assess** risk by weighing two things: the **likelihood** 可能性 of an attack and the **severity** 严重性 of the damage.

Likelihood itself depends on the **value** of the target (adversaries chase what looks worth stealing), the **skill** needed to exploit the vulnerability (a well-documented exploit needs little skill, so more adversaries can use it), and the **motivation and capability** of likely adversaries. Severity is usually measured in **financial** cost, but includes **reputational** and **operational** damage too.

The final rating can be written two ways, and the exam wants you to tell them apart:

- **quantitative** 定量 - a number: a score on a scale (e.g. 1-10), or a money value (e.g. "a \$10,000 annual risk").
- **qualitative** 定性 - a label: *low / medium / high / severe*, or a grid such as *likely-high-impact* vs *unlikely-low-impact*.

A written **risk assessment** 风险评估 should record, for each risk: the vulnerable asset and its value, the likely threats, how the specific vulnerability would be exploited, the severity if it were compromised, and a final quantitative or qualitative rating.

Once a risk is measured, an organisation has four ways to **manage** it:

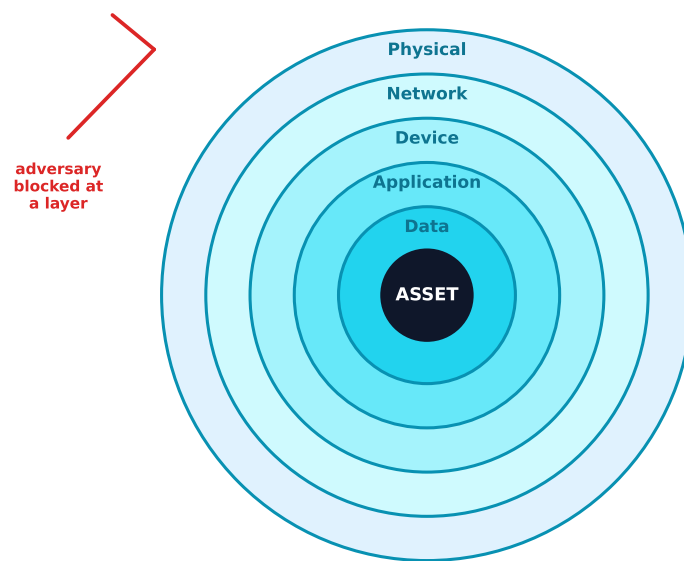
- **Avoid** 规避 - stop the risky activity (only possible if it isn't essential).
- **Transfer** 转移 - shift the burden to someone else, such as an insurer.
- **Mitigate** 缓解 - add controls to lower the likelihood or impact.
- **Accept** 接受 - live with the leftover **residual risk** 剩余风险, because perfect security is impossible.

Security controls are grouped two ways. By **type**: **physical** 物理 (locks, fences, guards), **technical** 技术 (firewalls, anti-malware, encryption), and **managerial** 管理 (policies and procedures). By **function**: **preventative** 预防性 (stop an attack, like a lock),

detective 检测性 (spot an attack, like a camera), and **corrective** 纠正性 (fix and restore, like patching).

Worked example. A hospital stores patient records on an unencrypted server in an unlocked room. Rate the risk: the asset is highly sensitive (patient data, protected by law) *and* the vulnerability is easy to exploit (no encryption, no access control), so this is a **high** risk. Now classify one fix - a door lock: by **type** it is a *physical* control, and by **function** it is *preventative* (it stops entry before an attack even begins).

The best strategy layers many controls - a **defense-in-depth** 纵深防御 approach. If an adversary bypasses one layer, another still stands. Layers include human, physical, network, device, application, and data.



Defense in depth: many layers so one breach does not expose the asset

Physical Vulnerabilities and Attacks

Digital security means nothing if an adversary can simply walk in. Common **physical attacks** 物理攻击 often begin with social engineering:

- **Piggybacking** 尾随 (获许可) - tricking an authorised person into holding a door open (for example, by carrying a heavy box).
- **Tailgating** 尾随 (未察觉) - slipping through a secured door behind someone without their knowledge.
- **Shoulder surfing** 肩窥 - watching someone type a password or read sensitive information.
- **Dumpster diving** 翻垃圾搜集情报 - searching a target's trash for useful information.
- **Card cloning** 门禁卡复制 - copying an access card to enter restricted areas.

With physical access, an adversary can cut power, steal or copy data, or plug in a **keylogger** 键盘记录器. We rate physical risk as **high** when sensitive systems sit in a space without controlled access, **moderate** when an unimportant area could act as a **foothold**

立足点 to reach other resources, and **low** when the asset is worthless and unlikely to be attacked.

Protecting Physical Spaces

Managerial controls come first: **security-awareness training** teaches staff not to badge strangers in, and a **workstation security policy** requires locking devices, clearing desks (a **clean desk policy** 清桌政策), and using privacy screens.

Physical controls then harden the building: fences, gates, and **bollards** 防撞柱 deter access; locks protect doors and cabinets; **card readers** 读卡器 log and restrict entry; an **access control vestibule** 门禁前室 (a two-door airlock) stops piggybacking; disabling **USB ports** blocks malware drives; and an **uninterruptible power supply (UPS)** 不间断电源 keeps devices running through an outage. Organisations prioritise these by matching the cost of a control to the severity of the risk.

Detecting Physical Attacks

Some controls **detect** attacks rather than prevent them. **Cameras** record activity and help after-incident investigations; **security guards** respond to what they see; **motion sensors** 运动传感器 alert staff to movement; and employees themselves often notice an intruder first.

Placement matters. Cameras belong at **points of ingress and egress** 出入口 (entrances and exits). Motion sensors work best in low-traffic areas like server rooms - put them in a busy hallway and constant false alarms make everyone ignore them. Stationary guards protect a fixed high-value point, while patrolling guards are harder for an adversary to plan around. Reviewing **door-open times** in entry logs can even reveal piggybacking, because a door held open too long is suspicious.

Exam tips

- Memorise the **CIA triad** and be ready to say *which* goal a control protects - encryption serves **confidentiality**, a hash checks **integrity**, a backup restores **availability**.
- Know the **four risk responses** (avoid, transfer, mitigate, accept) and the two ways to **classify controls** (by type: physical/technical/managerial; by function: preventative/detective/corrective).
- Distinguish **piggybacking** (with consent, tricked) from **tailgating** (without the person's knowledge) - exam questions test this exact pair.
- For risk-rating questions, **high risk needs both high value AND easy exploitation**; a "foothold to other systems" is the classic **moderate** risk.
- **Defense in depth** is the model answer whenever a question asks why one control is not enough.