

华为乾崑智能汽车解决方案 网络安全白皮书



深圳引望智能技术有限公司

深圳市龙华区福城街道九龙山社区景悦路1号A1栋101 邮编：518110

<https://auto.huawei.com/cn/>

免责声明

您购买的产品、服务或特性等应受深圳引望智能技术有限公司商业合同和条款的约束，本文档中描述的全部或部分产品、服务或特性可能不在您的购买或使用范围之内。除非合同另有约定，深圳引望智能技术有限公司对本文档内容不做任何明示或暗示的声明或保证。

由于产品版本升级或其他原因，本文档内容会不定期进行更新。除非另有约定，本文档仅作为使用指导，本文档中的所有陈述、信息和建议不构成任何明示或暗示的担保。

版权所有 © 深圳引望智能技术有限公司2025。保留一切权利。

未经本公司书面许可，任何单位和个人不得擅自摘抄、复制本档内容的部分或全部，并不得以任何形式传播。

联合发布单位（排名不分先后）

引望

SEES 赛力斯

BAIC 北汽新能源

阿维塔

长安汽车
CHANGAN

江汽集团
JAC GROUP

东风汽车
DONGFENG MOTOR

广汽集团
GAC GROUP

SAIC 上汽集团
SAIC MOTOR

CATARC 中汽智能科技(天津)有限公司
CATARC Intelligent Technology(Tianjin)Co., Ltd.

中国汽研
CAERI

上海机动车检测认证
技术研究中心有限公司
Shanghai Motor Vehicle Inspection Certification
& Tech Innovation Center Co., Ltd.

Contents

华为乾崮智能汽车解决方案
网络安全白皮书

目录

	前言	01
PART 1	华为乾崮智能汽车解决方案概述	05
	1.1 华为乾崮智能汽车解决方案简介	06
	1.2 华为乾崮智能汽车解决方案技术特征概述	08
PART 2	华为乾崮智能汽车解决方案网络安全理念与架构	11
	2.1 零信任网络理念	12
	2.2 纵深韧性架构设计模型	13
	2.3 网络安全解决方案架构	15
	2.4 产品网络安全基线	16
	2.4.1 网络安全基线制定的原则	17
	2.4.2 产品网络安全基线介绍	18
	2.5 AI(Artificial Intelligence)安全	21
	2.5.1 AI安全面临的挑战	21
	2.5.2 构建安全可信的AI系统	26
PART 3	华为乾崮智能汽车解决方案网络安全技术	30
	3.1 基础安全	31
	3.1.1 系统安全	31
	3.1.2 数据安全保护	33
	3.1.3 可信互联通信	35
	3.2 安全解决方案	36
	3.2.1 远控端到端保护技术	36
	3.2.2 一车一授权防刷机技术	37



前言

智能汽车正在重塑人类的出行体验，但同时也面临着前所未有的网络安全挑战。随着汽车从机械产品向"移动的数据中心"转变，网络安全已成为智能汽车的重要因素。每辆现代智能汽车搭载着上百个ECU（电子控制器）、运行着数亿行代码、每天产生TB级数据，网络安全漏洞可能带来严重影响。

智能汽车面临的关键网络安全挑战：

- 威胁持续演进：从简单的车辆盗窃到复杂的数据窃取、隐私侵犯、甚至人身安全威胁；
- 攻击面急剧扩大：从传统的物理接触攻击扩展到远程网络攻击，攻击数量快速增长；
- 法规要求趋严：UN R155/R156、GDPR（通用数据保护条例）等全球法规，《网络安全法》《数据安全法》《个人信息保护法》《汽车数据安全若干管理规定（试行）》、GB 44495-2024《汽车整车信息安全技术要求》、GB/T 44464-2024《汽车数据通用要求》等中国法规或标准，对汽车网络安全提出严格要求。

网络安全威胁演进态势：

- 1、汽车网络安全威胁呈现出专业化、规模化、产业化的特点；
- 2、攻击技术日趋高级化和复杂化：从简单的CAN消息注入发展到复杂的多阶段APT攻击；
- 3、攻击目标从车辆扩展到数据：攻击者更关注有价值的用户数据和商业机密；
- 4、攻击规模从单车扩展到批量汽车：针对车联网平台的攻击可能影响数百万辆车；
- 5、攻击动机多元化：除了经济利益，还包括政治目的、恐怖主义等。



3.2.3 高安全虚拟机技术	38
3.2.4 信任环车内通信安全技术	39
3.2.5 诊断安全技术	40

PART 4	华为乾崮智能汽车解决方案 网络安全法规标准遵从、测评与认证	42
4.1	GB 44495遵从	43
4.2	UN R155、ISO/SAE 21434遵从	44
4.3	C-ICAP测评	45
4.4	IVISTA测评	46
4.5	CC EAL安全认证	48
4.6	车联网云平台CCRC认证	49
4.7	汽车隐私保护测评规范	50

PART 5	华为乾崮智能汽车解决方案 全生命周期安全保障	51
5.1	安全隐私可信工程	52
5.2	VCSL网络安全实验室	54
5.3	漏洞管理	55
5.3.1	漏洞管理是保障网络空间安全的重要工作	55
5.3.2	漏洞管理是产业链的共同责任，需要协同合作来共同解决	56
5.3.3	引望遵循行业标准和最佳实践建立漏洞管理体系和工程能力	57
5.3.3.1	漏洞管理目标	57
5.3.3.2	漏洞管理原则	58
5.3.3.3	漏洞处置关键阶段	59
5.3.3.4	漏洞管理行为准则	59
5.3.4	引望漏洞管理框架与实践	61
5.3.4.1	全量漏洞管理	61
5.3.4.2	全生命周期漏洞管理	62
5.3.4.3	全供应链管理	63
5.3.4.4	漏洞管理平台	64
5.4	应急响应	64

A 缩略语表 / Acronyms and Abbreviations	66
--	-----------



攻击面持续扩大，按照攻击危害和攻击距离做如下分类：

1、按攻击的危害程度分类

智能汽车作为信息系统和载人工具的结合，尤其关心对行车状态的影响。按受到外部攻击后的危害程度可以分为以下三种：

1）控制车辆行驶状态：此类控制可以直接威胁人身安全，控制的功能包括：启动、停止、加速、转向、制动、挂挡等。

2）控制车身状态或者车辆状态配置：此类攻击影响到车辆的正常使用，可以间接影响人身安全，但在及时、有效干预下，可以避免人身安全威胁。被控制的功能包括：门、窗、车尾箱、座椅、空调、雨刮器、灯、喇叭、音响、仪表盘、悬挂、摄像头、激光雷达、驾驶模式设置、充电设置等。

3）不影响车辆状态的攻击：此类攻击影响车辆信息系统的功能，影响数字资产、隐私数据、财务等，包括：音视频窃听、位置跟踪、支付账号盗取、照片录像盗取等。此类攻击不影响人身安全。

2、按实施攻击的距离分类

对汽车的攻击有多种分类，其中，攻击者与汽车的距离，是在汽车网络安全中关注度比较高的一种分类，依次可分为远距离攻击、近距离攻击、物理接触攻击（无距离攻击）。

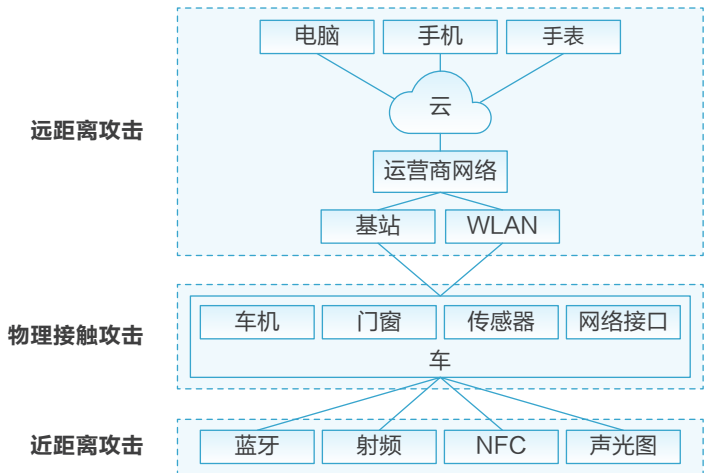


图1：按实施攻击的距离分类



1）远距离攻击：攻击者与汽车的距离没有限制，通过互联网实施无限距离的攻击。攻击者获得云的控制权，是远距离批量控制汽车的关键。

2）近距离攻击：攻击者在汽车附近出现，通过近距离通信方式攻击汽车。攻击方式包括蓝牙、射频（如射频钥匙）、NFC(如NFC钥匙)、声音（如超声波、智慧语音控制）、光线图像（如激光照射攻击）等。此类攻击需要在汽车周边约200米以内实施，可以不进入汽车座舱内，不破坏汽车门窗，实施无物理接触的攻击。

3）物理接触攻击（无距离攻击）：攻击者通过硬件连线接入汽车实施攻击，包括：a）进入座舱内部，连接诊断口、USB口，点击屏幕进行操作；b）打开电门盖，接入充放电口；c）破坏车灯或者车壳，获得内部连线的接口。

4）组合攻击：组合运用上述多种攻击手段达成攻击目的。此类组合攻击可以实现一次物理接触后，持久化实现随时的远距离控制。

全球各国政府高度重视汽车网络安全，纷纷出台相关法规、标准或要求：

欧盟：UN R155/R156法规要求汽车制造商建立网络安全管理体系（CSMS）。

美国：NHTSA 发布了《现代汽车的网络安全最佳实践》。

中国：制定《网络安全法》《数据安全法》《个人信息保护法》《汽车数据安全若干管理规定（试行）》，发布GB 44495-2024《汽车整车信息安全技术要求》，形成了完整的监管框架。

日本：制定《自动驾驶汽车安全技术指南》，明确网络安全要求。

合规挑战主要体现在：法规要求复杂多样、标准更新频繁、跨国合规成本高昂、违规处罚日趋严厉。企业须建立全球化的合规管理体系，才能在各个市场顺利开展业务。

为应对上述智能汽车的网络安全挑战，华为乾崮智能汽车解决方案提出了一套分层纵深防御的安全架构，以整车安全为目标，跨域协同，为用户提供全方位的网络安全保护，以下称为华为乾崮网络安全解决方案。



白皮书主要观点：

- 1、智能汽车网络安全需要体系化思维，不是单点技术问题；
- 2、安全必须从设计阶段开始融入，而非事后补救；
- 3、持续运营和演进是保持安全有效性的关键；
- 4、技术创新与合规要求必须并重发展。

本文介绍华为乾崑智能汽车解决方案（以下简称“华为乾崑”）的网络安全理念与实践。用户可以通过本文了解华为乾崑提供的全方位网络安全保护。

本文主要从以下章节进行阐述：

- ▶ 前言：简要说明业界智能汽车面临的网络安全威胁，以及华为乾崑网络安全的业务背景。
- ▶ 第一章：华为乾崑网络安全概述，介绍智能汽车的典型电子电器架构、技术特征、与网络安全相关的业务。
- ▶ 第二章：华为乾崑网络安全的理念与架构，介绍网络安全解决方案设计所采用的零信任网络理念，基于ISO/SAE 21434标准对整车进行安全威胁分析，提出分层纵深的网络安全架构。
- ▶ 第三章：展开介绍重要的华为乾崑网络安全解决方案，针对不同的网络安全风险，实施全方位的网络安全防护技术方案，包括基础安全、远程端到端加密、诊断安全、信任环车内通信安全、一车一授权防刷车等，联合构成整车安全防线。
- ▶ 第四章：介绍华为乾崑网络安全高质量符合国家标准、国际标准，在行业推荐性安全标准测评认证中取得优异成绩。
- ▶ 第五章：介绍华为乾崑网络安全如何构建全生命周期安全，保证过程安全、结果安全。通过优秀的安全工程实践保证安全质量，对智能汽车进行安全测试，在车辆上市前保证问题闭环；在车辆上市后，时刻关注业界安全动态，第一时间进行安全应急响应，确保汽车的全生命周期安全。

PART 1

华为乾崑智能汽车解决方案概述

深圳引望智能技术有限公司（简称“引望”，下同）坚信“质量为企业的生命”，坚持“以质取胜”。我们致力于为客户提供高质量的产品和服务；我们将多年的质量管理经验延伸到所有的合作伙伴，构建全流程、端到端的质量管理体系，共同推动相关产业链的高质量发展。网络安全和隐私保护是数字化、智能化世界发展的基石。作为安全可信有能力、开放合作负责任的智能汽车解决方案提供商，引望坚持将对网络和业务安全性保障的责任置于公司的商业利益之上。为此，引望持续构建完善的网络安全治理体系，在端到端的业务流程中嵌入网络安全要求，协同伙伴提供安全可信的产品、解决方案和服务，支持客户增强网络韧性。同时，我们积极为网络安全标准贡献力量，与利益相关方一同，合作共建安全可信的网络空间。



1.1 华为乾崮智能汽车解决方案简介

华为乾崮智能汽车解决方案，包括乾崮智驾、鸿蒙座舱、乾崮车控、乾崮车载光、乾崮车云服务，并以此为核心，实施一系列整车跨域业务解决方案，实现电动化、网联化、智能化、共享化，为用户提供智能、舒适、安全可靠的驾乘体验。

► **乾崮智驾——引领产业迈入新时代：**ADS（辅助驾驶系统）全新端到端架构、全向防碰撞系统、车位到车位和泊车代驾等业界领先功能，人驾更安全，智驾更舒心。

► **鸿蒙座舱——智能出行空间：**基于HarmonyOS开发的鸿蒙座舱重新定义智能座舱体验，更懂你的小艺语音，多人、多屏、多音区联动，丰富的车机应用和非凡的车载音响，让用户操作得心应手，让体验丰富多彩。

► **乾崮车载光——智能汽车新视界：**颠覆视觉体验，打造智能车载光解决方案，提供AR-HUD、百万像素车灯模组、光场屏等产品。助力汽车智能化创新，让驾驶更安全，交互更便捷，驾乘更舒适。

► **乾崮车控——数据底盘引擎：**内置自研车控模组，实现6合1全域融合架构，实现中央集中全维协同控制，带来极致驾乘体验。

► **乾崮车云服务——全生命周期极致体验和贴心服务：**提供全系手机无感解闭锁的数字钥匙服务；实时查看车周情况的乾崮云瞰服务；远程自定义迎宾灯语的乾崮云曦服务；一路记录的行车记录仪；可随时控车的远程控制功能。

华为乾崮为智能车解决方案提供了一套智能、互联、安全可靠的系统，也对网络安全保护提出了更高的要求。安全风险较高的功能有：



► **随时联网：**智能汽车通过蜂窝网和WLAN随时连接互联网，并可以安装APP应用，在提供智能、丰富体验的同时，也增加了远程攻击与数据泄露的风险。

► **远程控制：**智能汽车提供通过手机、手表远程控制汽车的功能，包括开门开窗、远程启动、远程召唤等功能，带来非授权远程控制汽车的风险。

► **无感解闭锁：**智能汽车提供数字车钥匙功能，通过手机蓝牙车钥匙，无需用户操作即可解锁车辆，带来被攻击无感开门的风险。

► **辅助驾驶：**辅助驾驶系统具备辅助控制动力、转向、刹车等功能，给攻击者提供了软件控制动力的可能性，带来远程攻击控制动力域的风险。



1.2 华为乾崮智能汽车解决方案技术特征概述

华为乾崮在智能网联方面有卓越的进步，提供丰富的网联功能，并且具备辅助驾驶能力。智能汽车对外提供丰富的互联接口，内部具备丰富的传感器和辅助驾驶控制能力。整车架构由外到内，可以分为车云与手机、车外通信层、车内通信层、智能控车层、执行ECU等几个层。

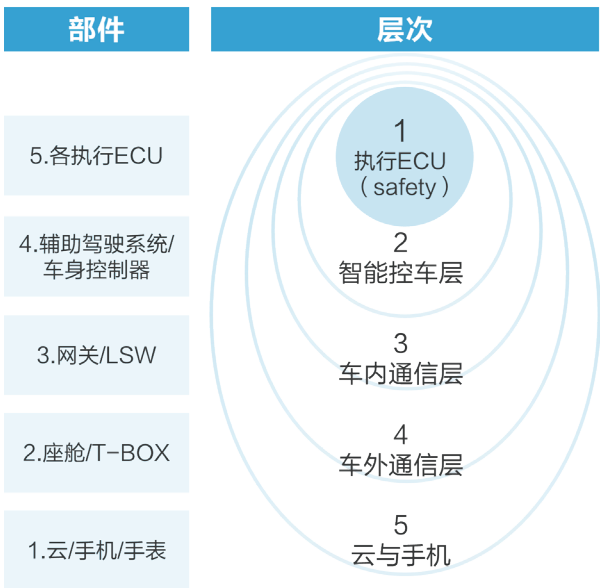


图2：安全架构分层

1、云与手机

1) 云：包括车云、ADS云、智能车连接的其他云。车云为智能车提供丰富的在线服务，包括在线升级、远程控制、智慧寻车、哨兵模式等。ADS云为辅助驾驶系统提供地图更新、路线规划等功能。智能车上的APP会连接其服务器所在的云，进行账号登录、数据同步、服务内容交互等。

2) 手机和手表：手机上安装有车主APP，配合手机系统的能力，提供远程控制、数字钥匙、遥控泊车等功能。智能手表和运动手表也能与车辆进行交互，实现远程控制、数字钥匙等功能。

2、车外通信层

1) 智能座舱：智能座舱是指车内驾驶员、乘客乘坐的空间，包括智能车机、氛围灯、音响、抬头显示等软硬件。其中，智能车机是智能座舱的核心部件，提供人机交互、车身控制、音视频娱乐等功能。车机对外进行云连接、人机交互、WLAN和蓝牙通信，包含有麦克风、摄像头和USB接口，所以归属在车外通信层。车机对内可以设置门窗、空调、座椅等，同时车机内置有导航、行车记录仪、哨兵模式、通话等功能，并可以通过应用市场安装第三方应用，提供音乐、电影、短视频等功能。

2) T-BOX：为智能汽车提供远程网络连接的功能，可通过蜂窝网连接上网。远程控制指令通过云端下发到T-BOX，再转到车内执行。T-BOX还提供数字车钥匙功能，可自动连接手机，提供无感解锁车门、手机遥控泊车等功能。

3、车内通信层

1) 网关：为整车以太网和CAN总线提供连接的网络中心设备。除了提供连接功能以外，网关具备车辆功能管理、车身控制能力，例如OTA升级管理、诊断管理、VHR日志管理、证书密钥管理等。网关对外提供诊断接口，用于连接诊断仪进行车辆检修，可以读故障码、升级系统、配置汽车参数。

2) 局域网交换模块：即LAN Switch模块，可内置于SOC（System-On-a-Chip，片上系统）中，也可能作为独立芯片，支持车内以太网报文转发。

4、智能控车层

1) 辅助驾驶系统：实现城区辅助驾驶、高速辅助驾驶、远程召唤、辅助自动泊车、遥控泊车的功能，该系统从智能座舱或者手机获得用户的指令，从激光雷达、毫米波雷达、摄像头等传感器获得外部道路信息，依据地图规划行进路线，控制车辆的动力系统、转向系统、制动系统，辅助控制车辆行驶到目的位置。

2) 车身控制器：车身控制器提供对门、窗、座椅、空调、灯光等车身设备的控制功能。车上有多个车身控制器，例如左侧、右侧、前侧各一个，车上各设备可以就近接入车身控制器。车身控制器通过以太网连接车机、T-BOX，接收控车指令，通过CAN和LINE连接，对下挂设备进行控制。

5、执行ECU

域控制器是汽车功能域（如动力、底盘、车身、智能座舱、辅助驾驶）的集中运算单元，由硬件（主控芯片、存储器件等）和软件（操作系统、中间件、应用算法）构成，实现对域内功能的统一调度与控制。包含辅助驾驶域控制器、智能座舱域控制器、车身域控制器、联网域控制器等。除了上述域控制器以外，整车还包含很多其他ECU，此类ECU负责单项功能的执行，统称为执行ECU。典型的执行ECU有：电池管理系统、热管理系统、电机、灯光、音响、AR-HUD（增强现实平视显示器）、激光大灯、投影屏等。

华为乾崑智能汽车解决方案
网络安全理念与架构



2.1 零信任网络理念

2010年，约翰·金德维格（John Kindervag）首次明确提出“零信任”这一术语，他认为零信任本质是以身份为基石的动态访问控制。该技术强调以身份为核心，运用动态访问控制技术，将细粒度的应用、接口、数据作为重点保护对象，并遵循最小权限原则，以此构筑端到端的安全边界。这一阐述为零信任理念奠定了理论基础，使其从模糊设想走向清晰概念。

2014年，国际云安全联盟在零信任理念基础上提出软件定义边界（SDP），并发布标准规范，进一步推动零信任从理论走向落地应用。SDP协议致力于打造按需、动态配置的隔离网络环境，将受信任网络与不安全网络有效隔离，缓解基于网络的攻击。

零信任网络摒弃了传统网络安全中“内网默认可信”的假设，主张对网络中的任何访问请求，无论来自内网还是外网，都不应给予默认信任。它以身份为中心，围绕用户、设备、应用等构建动态访问控制机制，将安全防护的边界从网络边界延伸到每一次访问行为，确保只有经过授权且安全的访问才能被允许，从而有效抵御来自内部和外部的网络威胁。

关键原则

最小权限原则：仅授予用户或设备完成特定任务所需的最小权限，且权限会根据环境变化动态调整。

持续验证与评估：在访问过程中，持续监控设备状态和网络环境。一旦发现如设备出现未授权软件等异常，立即重新评估访问权限，甚至阻断访问。

零信任网络通过身份认证、设备安全评估、访问策略管理和持续监控等技术模块协同工作。身份认证模块负责核实用户身份；设备安全评估模块检测设备是否存在漏洞、是否被恶意软件感染等；访问策略管理模块依据身份和设备状态制定并执行访问策略；持续监控模块实时收集和分析网络数据，发现潜在威胁并及时响应。



将零信任网络理念用于智能汽车网络安全架构设计，应摒弃“内网默认可信”的假设，车内不同部件不再视为同一个信任网络，在部件之间的访问需要采取接入认证、访问控制等安全措施。但车内ECU多达100多个，无法在每两个ECU之间进行访问控制。在华为乾崮网络安全解决方案中，将部件进行分层部署，以执行ECU为被保护层，构建多层纵深的保护架构。采取零信任网络理念进行设计，外层访问内层需要进行接入认证、权限鉴定，保护访问的完整性、机密性、可用性，做到防篡改、抗仿冒、防重放。

2.2 纵深韧性架构设计模型

依据零信任网络的理念，对车内部件进行分层分域设计，各层间进行访问控制。依据车内部件功能的划分，可以将车内设备分为车外通信层、车内通信层和智能控车层。车外通信层对智能控车层的访问，进行接入认证、访问控制等安全保护，设置隔离机制。

依据ISO21434的威胁评估方法，智能控车层是对车辆功能影响很大的分层，应该放置在内层进行保护，使得攻击可行性降低。车外通信层从业务上需要对外暴露接口，从布局的暴露面维度看，是攻击可行性相对高的分层，应该尽量剥离控车功能，使得攻击影响程度最低化，以此提高整车的抗攻击能力和韧性。

风险矩阵		攻击可行性（仅考虑暴露程度）			
		非常低	低	中	高
影响程度	严重	2 辅助驾驶系统网关	3	4 避免本区域	5
	重要	1	2	3 智能座舱	3
	中等	1 执行ECU	2	2 TBOX	2
	可忽略的	1	1	1	1



在架构设计中，避免出现攻击可行性高且影响程度严重的部件。例如，辅助驾驶系统和网关，影响程度高，在架构布局中尽量降低其攻击可行性，布局在电子电器架构的内层；智能座舱和T-BOX，业务上需要联网和开放外部接口，攻击可行性相对高，则尽量降低其影响程度，限制其对动力和底盘的直接访问，对此类部件控车进行接入认证和访问控制。

按照此理念进行设计，得到以下图示的架构模型。该模型作为各车型电子电气架构安全设计的指导，不代表各款车型都完全遵从本架构实施，不同车型的架构依据实际功能会有所不同。

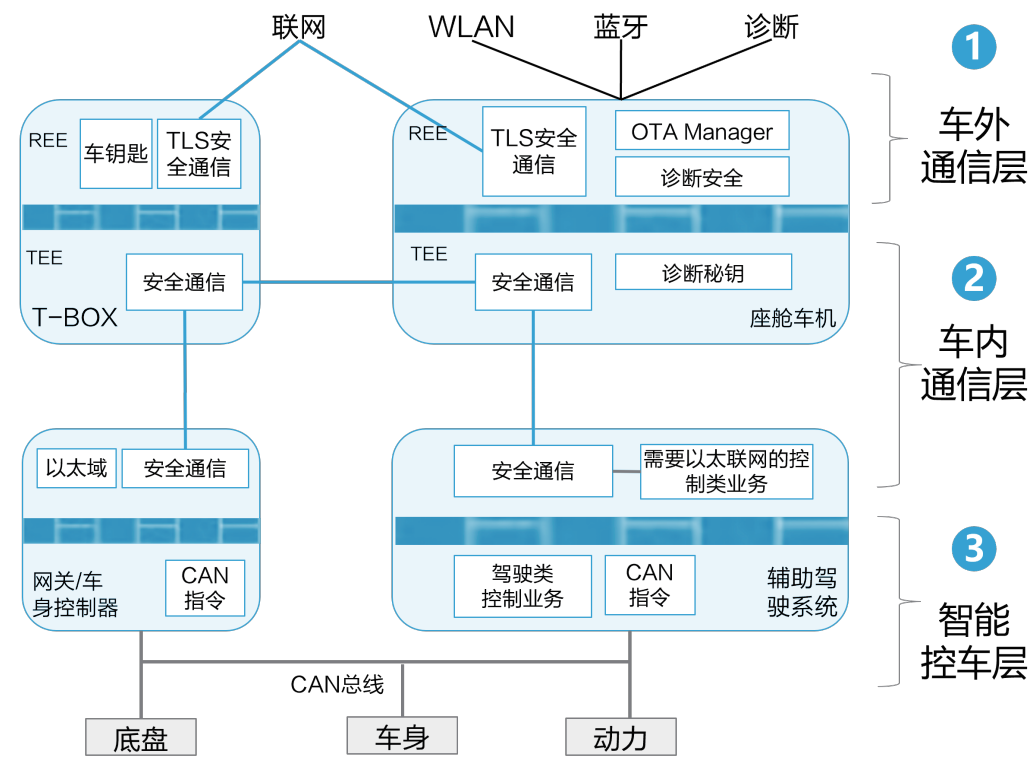


图3：分层分域设计架构模型

1、车外通信层：整车对外接口规整到座舱车机和T-BOX，系统做充分的防御加固

- 1) 联网、蓝牙、WLAN、诊断接入到座舱车机和T-BOX的REE（Rich Execution Environment，富执行环境）域，完成接入认证和指令、数据校验。
- 2) 用TEE（Trusted Execution Environment，可信执行环境）隔离车内通信功能，使得REE被攻击不易渗透到车内。

2、车内通信层：控车指令使用安全通信通道，实现以太通信的完整性、机密性保护

- 1) 以TEE为车内安全通信的起点，安全通信密钥保存在TEE中。攻破REE不易发起安全通信。
- 2) TEE对REE的访问进行权限控制，防止攻击者访问通信密钥。

3、智能控车层：包括网关、车身控制器、辅助驾驶系统，将以太网连接与CAN总线连接做隔离，基于业务做防火墙，保护车内的CAN总线和执行器ECU。

- 1) 以太网连接和CAN连接隔离，一般以太网连接由SOC提供，CAN连接由MCU提供。MCU具备下发CAN指令的权限，可直接控制动力、底盘。MCU对SOC采取零信任的策略，对SOC进行访问控制和指令过滤。
- 2) 智能座舱和T-BOX设备不能直接访问CAN总线，防止直接影响车辆行驶功能，降低攻击的影响程度。

2.3 网络安全解决方案架构

在华为乾崮网络安全解决方案中，基于模型，设计分层的网络安全架构，逐层进行零信任的安全防御。在各层实施一系列网络安全跨域方案，协同构建成整车的纵深防御解决方案。



图4：分层的网络安全架构

1、远程控车实现端到端保护，云端不能控制车辆，有效防御批量控车

安全思路：由终端签名，由车端验签，只有车主终端发出的指令才能控车，有效防御非授权控车。

2、落地一车一授权，软件版本与芯片绑定，专车专用，保障系统完整性

安全思路：即使OTA（Over-the-Air Update，空中下载更新）包泄漏、被替换，也不能刷写到车。

3、落地高安全虚拟机，隔离娱乐域和仪表域，虚拟机达到CC EAL6+认证

安全思路：签发控车指令的TEE入口实现全路径权限最小化。

4、落地车载信任环，智能部件实现P2P（点对点）证书认证，私钥不出硬件，专车专用

安全思路：只有合法设备、TEE签名的合法指令才能跨部件执行。

5、实现诊断仪由云端统一授权管控，一车一件一密，防止通过OBD攻击

安全思路：只有具有合法证书、云端授权的诊断仪才能接入，阻断非法攻击。

2.4 产品网络安全基线

引望将网络安全管理要求融入到内部各业务流程中。在集成产品开发IPD流程中，将产品网络安全要求融入到从规划、设计、开发、验证、版本发布到全生命周期管理的整个过程中，保障了引望交付给客户的全系列产品和版本都能满足各利益相关方关注的安全质量要求。在这些管理要求、工程与技术规范中，引望产品安全基线是强制性的基础要求。基于过去十多年在产品安全质量上的长期实践,我们持续不断的优化、刷新这一产品安全基线，也和合作伙伴、供应商等共同分享这个安全基线对保障端到端供应链安全所带来的价值。安全实践也证明了，产品安全基线是一个有效管理产品安全质量的方法，也保障了产品在客户网络上一直以来良好的安全运行记录。

2.4.1 网络安全基线制定的原则

我们认为：网络安全基线置于商业利益之上

- 面对安全事件频发的网络安全现状，提升产品自身安全是关键的风险削减措施之一。
- 将网络安全管理嵌入到产品开发过程中，使网络安全成为产品的基础能力之一，才能从根本上解决网络安全问题。
- 提炼基础的、通用的产品安全基线要求，并在所有产品中落地，可以确保所有产品都能达到一致的安全质量基础要求，并随着基线的刷新而不断提升产品安全质量。
- 在端到端的网络安全框架下，产品安全基线作为基本的安全要求，嵌入产品开发流程，通过严格的质量保证活动，确保产品安全质量，减少发生安全事故的风险。

产品安全基线原则：

- 以“基于结果、普适、无差别、持续优化”为原则，制订引望产品安全基线，保证了产品安全基线的有效落地、可验证和可持续提升产品安全质量水平。
- 通过研究适用法律法规，深入理解法律合规要求、客户业务诉求、业界最佳实践、行业内已知问题等，识别出产品通用的、最关键的安全要求，制定了产品安全基线。

为了保证引望产品安全基线的有效落地、可验证、可持续提升产品安全质量水平，我们在制定和管理产品安全基线遵循以下原则：

► 它是基于结果的

只有结果性的要求才真正是客观的、可验证的，基于结果的产品安全要求有助于产品准确地实现预期目标，也有助于客户基于客观结果比较和选择所需的产品。

► 它是普遍适用的

进入产品安全基线的要求须是通用的、普适的，能够适用所有或者绝大部分产品、适用于多种业务场景，以确保所有产品都能达到一致的安全基础质量要求。



► 它是无差别的

我们处在全球供应链的时代，如果我们要真正确保所有系统和所有基础设施的安全，就应该对所有的供应商(无论在哪里开发、或者在哪里生产)、所有的产品部件(无论产品被应用于哪个区域、哪个网络、或者网络的哪个位置)采取同样的安全基础标准。

► 它是持续优化的

网络安全是一个动态的过程，产品安全基线须持续定期刷新，以适应不断发展的网络安全形势。

2.4.2 产品网络安全基线介绍

引望通过对法律法规、标准规范、安全实践等的分析总结，提炼产品共同的、最关键的安全要求，制定了产品安全基线，并随着网络安全和业务的发展持续刷新。

产品网络安全基线			
1.保护用户通信内容	4条	9.敏感数据保护	4条
2.保护用户隐私	7条	10管理和维护安全	5条
3.防止后门	5条	11.安全启动和完整性保护	2条
4.防止恶意软件、恶意行为	2条	12.安全资料	3条
5.访问通道控制	5条	13.安全编码	2条
6.系统加固	4条	14.安全编译	1条
7.应用安全	3条	15.生命周期管理	2条
8.加密	5条	总计	112条

如上表，引望安全基线共计15类54条要求，为了便于不同应用场景准确理解和执行安全基线，同时还开发了实施指导与解读共计112条。



1、保护用户通信内容

用户通信内容受各国法律保护，在产品设计中严格遵循行业通用安全标准，确保通信数据安全。

2、保护用户隐私

安全基线分析了欧洲GDPR、德国、法国、英国、加拿大、中国等国家的隐私保护法律法规要求，参考GSMA（全球移动通信系统协会）隐私设计规范等业界实践，按照“合法、正当、透明”、“目的限制”、“数据最小化”、“准确性”、“存储期限最小化”、“完整性与保密性”、“可归责”等7条隐私保护基本原则用于指导产品设计和开发。

3、防止后门

安全基线不允许在产品中植入后门，也不允许其他人在华为乾崮设备中植入后门。

4、防止恶意软件、恶意行为

安全基线不允许产品存在病毒、木马等恶意软件，也不允许产品存在恶意广告、吸费、恶意消耗流量等恶意行为。

5、访问通道控制

采用隔离、认证等手段控制访问通道，能够有效降低攻击面，保证接入产品访问的安全。

6、系统加固

通过安全配置、补丁修复漏洞、删除或禁用不必要的服务等方法，对产品进行安全加强以提高产品的安全性和抗攻击能力。

7、应用安全

WEB应用、移动应用等各种应用容易受到黑客和恶意软件的攻击，导致未经授权的访问和修改，认证和鉴权是最基本的安全保护机制。

8、加密

密码算法是安全的基石，正确地使用密码算法是产品安全性的基石。根据不同应用场景使用密码算法，须使用公开的、经过专业评审和验证过的安全的密码算法，并正确地配置算法参数和选项。





9、敏感数据保护

在产品设计时，需要根据产品的使用场景，识别出产品中的敏感数据。典型的敏感数据包括认证凭据（如口令、私钥、动态令牌卡等）、加密密钥、敏感个人信息（如人脸、车辆位置等）等。从敏感数据的存储、传输和处理过程中采取认证、授权或加密等安全机制保证数据安全。

10、管理和维护安全

产品的管理维护采用严格的账号口令、安全的访问协议、完整的日志审计等安全机制，保证产品管理维护操作的安全。

11、安全启动和完整性保护

在产品的安装、升级中验证软件包的完整性，防止产品被攻击者恶意篡改。对于客户需要有高安全要求的产品同时需要考虑在启动过程中进行安全启动验证。

12、安全资料

产品的资料需要提供产品的通信矩阵、账号清单、安全加固和配置指南等安全资料，以便能够指导客户以最安全的方式使用、部署和维护产品。

13、安全编码

大部分安全漏洞很大程度上是由于代码编写的不规范、编程语言特性理解不当、接口调用不当等代码质量缺陷导致的，通过例行的静态代码安全和质量扫描，以及减少不安全函数的使用，可以提前识别代码中的安全质量缺陷，持续提升代码质量, 减少潜在的安全漏洞。

14、安全编译

编译器提供了相当多的安全选项来加固软件的安全性，通过添加安全编译选项，可以在软件中进行代码质量的安全加固，提高了攻击者的攻击难度，降低了代码质量缺陷变成可被利用漏洞的风险。

15、生命周期管理

产品软件需要使用安全记录良好并持续维护的开源和第三方组件，出现漏洞及时通过发布补丁或者升级到新版本等方式修复漏洞，确保产品软件在生命周期内的安全处于可维护状态。



2.5 AI (Artificial Intelligence) 安全



AI（人工智能）技术正在重塑汽车产业的未来图景。以辅助驾驶为代表，AI已深度渗透到辅助驾驶的感知、决策、控制全链路。华为乾崮智驾ADS 4的发布，为辅助驾驶技术发展带来新突破。其采用世界引擎 + 世界行为模型架构（WEWA架构），通过云端世界引擎（World Engine）运用难例扩散生成模型，可生成密度达真实世界1000倍的复杂场景数据，突破传统数据采集瓶颈；车端世界行为模型（World Action Model）基于全模态感知数据训练，大幅降低端到端时延，提升通行效率。鸿蒙座舱HarmonySpace 5采用全新混合大模型Agent架构（MoLA架构），整合通用大模型与音乐、影视等领域AI能力，构建丰富的软硬件生态，带来多空间场景体验。尽管AI驱动的创新持续改变汽车行业，但它们也带来了复杂的网络安全挑战，需要特别关注。

2.5.1 AI安全面临的挑战

从AI的发展过程可以看出，算法、算力、数据是AI发展的三个最主要的驱动力。而随着AI能力越来越强，引发人们对AI系统网络安全的广泛关注和思考。

1、安全合规与标准化挑战

整体而言，AI的技术储备、商业生态、文化与法律传统、以及所处的治理阶段等要素，都不同程度影响一个国家或地区的人工智能安全治理实践，因此，不同国家和地区的人工智能法律法规的制定和执行呈现一定的差异性。





欧盟：以风险为基础的分级监管

欧盟《人工智能法案》是全球首部全面且具影响力的人工智能法规，于2024年5月21日获批通过，并于2024年8月1日正式生效。这部号称史上“最严”的针对AI的全面监管法案，将AI产业链上的各主体均列入监管范畴，包括与欧盟市场有连接点的AI系统提供商、使用商、进口商、分销商和产品制造商等。需要注意的是，虽然这部法律已于2024年8月1日起正式生效，但部分条款的适用存在过渡期。根据欧盟委员会制定的时间表，法案中规定禁止的AI系统将在法案生效6个月后适用；通用型人工智能的相关义务和规则将法案生效后的12个月后适用；部分高风险AI系统的规则基于行业协调立法适配后将在36个月后将开始适用。

欧盟《人工智能法案》采用了以风险为导向的分级管理模式，将AI系统分为不可接受的风险、高风险、有限风险和最低风险四个级别。其中，认知行为操纵和社会评分等八大禁止类场景的AI系统被禁止进入欧盟市场；高风险AI系统（如自动驾驶、生物识别）需满足严格的监管要求，包括全面的风险评估、数据治理、透明度义务、人类监督等。此外，高风险AI系统的合规性评估贯穿其整个生命周期，包括设计研发、合格评估、系统注册和监测。欧盟还设立了域外效力，适用于所有在欧盟销售或部署AI系统的公司。如违反相关条款，企业最高可能被处以3500万欧元，或上一财年全球年营业总额7%的罚款（以数额较高者为准）。

标准方面，欧洲标准化组织CEN/CLC成立JTC21专门负责AI标准建设，该技术委员会主要通过引入ISO标准或独立开发的方式，系统性开展欧洲AI安全、可信标准建设。由JTC21发布、在研的AI安全相关标准主要包括：

- ▶ EN ISO/IEC 23894: 2024 Information technology – Artificial intelligence – Guidance on risk management
- ▶ EN ISO/IEC 42001: 2023 Information technology – Artificial intelligence – Management system
- ▶ prCEN/CLC/TR AI Risks – Check List for AI Risks Management，该标准还处于在研状态



美国：分散化与行业特定监管

美国联邦层面尚未出台系统性的人工智能安全法，采取分散的监管模式，由多个联邦机构和州政府共同负责。其重点在于强调AI系统的安全（safe）、可靠（secure）、可信（Trustworthy）。在2023年10月白宫发布的E.O.14110《关于安全、稳定和可信的人工智能行政令》中首次提出有硬性监管效力的大模型政府报告要求，要求“存在重大安全风险的强大双重用途基础模型开发者”，向政府报告并分享安全信息、测试信息等。但在2025年1月特朗普政府撤销了前总统拜登出台的第14110号行政令，否认了拜登行政令尝试构建的一个全面的AI治理框架，并在其随后发布的新行政令（E.O.14179）中表示愿意接受更高层次的潜在风险，呼吁各部门和机构修改或撤销根据拜登第14110号行政命令采取的所有政策、指令、法规、命令和其他行动，以推进美国在人工智能领域的领导地位。

为更好支持美国AI产业发展及国际化，NIST ITL实验室（美国国家标准与技术研究院信息技术实验室）将AI作为重点研究方向，围绕可信AI技术基础研究、AI标准化、AI技术创新等目标设定若干具体工作任务。NIST已正式发布的部分AI安全、风险管理相关标准主要包括：

- ▶ AI 100–1 Artificial Intelligence Risk Management Framework
- ▶ AI 100–2 Adversarial Machine Learning: A Taxonomy and Terminology of Attacks and Mitigations
- ▶ SP 800–218A Secure Software Development Practices for Generative AI and Dual–Use Foundation Models

美国NIST牵头开发的标准与实践，具备技术实用性，有助于开展AI安全管理及推动安全工程建设。

中国：发展与安全并重

2017年国务院《新一代人工智能发展规划》奠定了中国人工智能法治发展的总基调，提出“在大力发展人工智能的同时，必须高度重视可能带来的安全风险挑战，加强前瞻预防与约束引导，最大限度降低风险，确保人工智能安全、可靠、可控发展”的要求。《规划》在人工智能治理方面提出战略目标：一是到2025年，初步建立人工智能法律法规、伦理规范和政策体系，形成人工智能安全评估和管控能力；二是到2030年，建成更加完善的人工智能法律法规、伦理规范和政策体系。





目前，中国《网络安全法》《数据安全法》《个人信息保护法》等在涉及人工智能相关场景时可延伸适用。在规章层面，国家网信办联合有关部门，先后针对算法推荐服务、深度合成、生成式人工智能服务等不同应用领域进行针对性立法，保障人工智能健康发展、规范应用。

为落实在2023年10月18日发布的《全球人工智能治理倡议》，2024年9月9日全国网络安全标准化技术委员会（TC260）发布《人工智能安全治理框架》1.0版本，鼓励人工智能创新发展为第一要务，以有效防范化解人工智能安全风险为出发点和落脚点，提出了包容审慎、确保安全，风险导向、敏捷治理、技管结合、协同应对、开放合作、共治共享等人工智能安全治理的原则，并以此为基础，系统性开展人工智能安全标准体系建设。2025年9月15日国家网络安全宣传周网络安全技术高峰论坛上，《人工智能安全治理框架》2.0版等重要成果集中发布。

中国高度重视人工智能产业发展，标准正成为产业政策落地的重要抓手，全国网络安全标准化技术委员会(SAC/TC260)已制定并发布多项标准指导产业实践。包括：

- ▶ GB/T 45652-2025 网络安全技术 生成式人工智能预训练和优化训练数据安全规范
- ▶ GB/T 45654-2025 网络安全技术 生成式人工智能服务安全基本要求
- ▶ GB/T 45674-2025 网络安全技术 生成式人工智能数据标注安全规范

2、人工智能面临的内生安全挑战

人工智能技术的内生安全问题深植于其系统内部，涉及多个层面的安全挑战。从基础设施的薄弱性到数据处理的漏洞、从模型训练和部署的不规范到应用层和接口层的潜在攻击风险，这些安全隐患都可能威胁到AI系统的整体安全性与用户隐私安全。

• 数据安全风险

数据投毒

数据投毒是指攻击者在模型训练阶段集中注入恶意或经过精心设计的数据样本，使模型在训练后产生特定的错误行为或后门，导致模型在正常输入下表现正常，但在遇到特定触发条件时产生预期之外的行为。数据投毒对AI模型构成严重的安全风险，如在辅助驾驶应用中，可能导致严重的安全事故或经济损失。



数据隐私泄露

数据泄露是模型相关的敏感信息被非授权访问、使用或泄露。具体包括通过模型反向工程泄露训练数据，模型错误输出暴露个人身份等敏感信息，模型本身的结构和参数被窃取或滥用。

• 模型安全风险

对抗样本

对抗样本是针对模型的最典型威胁，它是一种通过对输入数据添加微小的扰动，使模型产生错误输出的技术。

对抗样本对模型的安全应用构成了严重威胁，其影响范围涵盖了多模态领域，包括视觉、听觉和文本处理系统。在视觉领域，辅助驾驶车辆可能误读被篡改的交通标志。听觉方面，语音助手可能被隐藏在背景音中的对抗指令操纵。人脸识别和生物特征识别系统面临被精心设计的图像或视频欺骗的风险。

提示词注入风险

提示词注入攻击指攻击者通过巧妙构造输入文本，来操纵AI模型执行非预期或潜在有害操作的一种恶意行为方法。

提示词注入对大模型构成了广泛的安全风险，涉及文本、图像、音频等多模态内容。这种恶意行为可能导致敏感信息泄露、系统执行非授权操作、生成误导性或有害内容，以及利用第三方应用漏洞进行欺诈。

• 应用安全风险

智能体安全风险

智能体（Agent）是一种能够感知环境、进行决策和执行动作的智能实体。随着大模型的能力增强，大模型可以充当智能体的大脑，通过规划、记忆、工具执行等组件完成各种复杂任务。大模型智能体的安全风险是指通过提示词注入或对抗样本的方式，让大模型规划出恶意的任务序列，或生成并执行恶意的指令。大模型智能体的安全风险可能导致严重的危害，例如造成实际的资金损失，泄露重要的个人信息等。大模型智能体往往具备访问或操作系统的敏感权限，会进一步加剧传统网络安全的风险。



算力底座安全风险

硬件层安全风险

利用侧信道技术，从硬件环境窃取关键模型信息。基于侧信道的模型窃取主要是在模型部署运行过程中通过操作系统或硬件等额外信息推断目标模型的机密属性，因此侧信道威胁的主要风险是模型属性推断，而一般恶意破坏者最感兴趣的模型属性就是目标模型架构信息。

操作系统层安全风险

如同其他类似的软件系统一样，AI系统运行也依赖于底层的操作系统、驱动等系统软件的支持。很多应用层软件不保护其模型，即使对于那些保护和加密模型的应用程序，攻击者能够通过简单的动态分析技术从应用层中提取模型。

三方件安全风险

AI已经演化成为了复杂的生态，AI系统中存在着大量的开源、第三方组件。恶意行为者完全可以通过找到组件未及时修补的漏洞，实现代码执行、模型窃取等恶意行为。

2.5.2 构建安全可信的AI系统

人工智能网络安全与隐私保护治理由公司成熟的网络安全与隐私保护组织体系负责。

1、人工智能网络安全与隐私保护治理原则

安全为本：将网络安全与隐私保护作为业务的基本要求来对待，并保证网络安全与隐私保护需求得到充分的资源投入。

合法合规：遵守国家法律法规和行业标准要求，确保AI业务的合法合规运行。融入业务：网络安全与隐私保护保障活动需要融入到各相关业务的政策、流程、规范、基线中，使之成为业务的基因，为客户提供安全可信的产品、解决方案和服务，保障业务成功。

融入业务：网络安全与隐私保护保障活动需要融入到各相关业务的政策、流程、规范、基线中，使之成为业务的基因，为客户提供安全可信的产品、解决方案和服务，保障业务成功。

主管担责：各级业务主管是所辖业务网络安全与隐私保护的第一责任人，各级流程责任人是所辖流程网络安全与隐私保护的第一责任人。

全员参与：所有员工具备网络安全与隐私保护意识和能力，在自身业务中落实网络安全与隐私保护基本要求。每个员工都要对自己所做的事情和产生的结果负责，不仅要对技术负责，也要承担法律责任。

独立验证：信任应基于事实，事实必须可验证，而验证必须基于共同标准。在此基础上，基于“不假定任何事情，不相信任何人，检验所有应检查的事项”理念，分层进行独立评估与验证。

开放合作：秉承开放透明的态度，真诚地与客户、供应商、合作伙伴、行业组织等利益相关方积极开展网络安全与隐私保护沟通与合作，责任共担、能力共建、价值共享，共同应对网络安全与隐私保护威胁与挑战。

持续优化：网络安全与隐私保护是一个持续的风险管理和能力建设过程，不存在绝对的安全，也不存在一劳永逸的方案，需要适时审视不足，持续改进网络安全与隐私保护管理和技术措施的适宜性、充分性和有效性。

在网络安全与隐私保护整体基调之下，结合业界标准和公司既有实践，采用治理加嵌入式管理的工作方法，以实现AI网络安全风险管控和合规合法的目标，构建人工智能网络安全治理的架构。



图5：人工智能网络安全治理架构



2、多层次防护护栏体系

由于AI自身的复杂性，仅仅通过单一手段难以有效控制安全风险，公司构建了多层次的防护护栏体系，确保AI系统安全可靠。

第一道护栏：数据安全

基于外部法规、标准和业务需求，构建面向辅助驾驶AI数据的全生命周期的过程保障能力，建立AI数据安全合规基线，打造AI数据安全合规流程及工程化能力，保障AI数据全生命周期的安全合规。

- ▶ 数据集来源检查：AI数据来源合法合规、采集可知可控，高风险数据集不进入数据生产线。
- ▶ 数据集内容检查：根据数据特征检查违规数据、隐私、版权风险。数据采集遵循最小化、本地化原则，非必要不采集，非必要不出车。对于业务确实需要出车的数据，应采取脱敏处理，并加密传输，防止数据泄露。
- ▶ 模型合规性检查：检查模型文件完整性的同时，检查模型对应数据集是否有合规处置记录和数据集追溯记录。

第二道护栏：模型安全

- ▶ 数据：持续构建安全数据集。通过多样化的数据集，利用监督微调来增强其模型的安全性。
- ▶ 算法：设计不同模态安全算法，旨在使AI系统能够避免已知的人工智能安全漏洞和恶意提示词攻击，从而确保模型在安全层面上实现对齐。
- ▶ 评测：采取人工与自动化相结合的方式安全测试。这种混合方法不仅提高了测试的效率，还增强了测试的全面性和准确性。

第三道护栏：模型安全风控

以明确的风险消减为目标，重点抓好输入端和输出端的风险把控，提供检测防御能力。如：内容不合规，隐私泄露、信息泄露，无法溯源，恶意代码执行等风险，如座舱AI大模型进行输出内容规范化检测，保证文明合规的语言环境，包括保护儿童用户。

第四道护栏：应用安全

第四道护栏以消减明确的智能体和应用框架风险为目标。首先，为了保障智能体正常运转，利用任务规划和任务执行护栏解决方案，保障任务规划和任务执行核心功能正常运转。其次，为了保障应用框架正常运转，利用API（Application Programming Interface）越权防护、访问控制、容器隔离、算子安全防护等技术，消减应用框架访问越权等风险。最后，为了保障应用安全必须加强传统的软件系统的全栈安全工程能力。



第五道护栏：运营运维安全

在AI系统运营和运维阶段，重点保护大模型系统所涉及的训练数据、模型、生成内容的全生命周期安全，构建大模型恶意为态势感知系统，提升拦截能力，对现网系统进行统一运营。通过对AI模型进行加密，防止模型泄露，减少被逆向以及漏洞攻击的风险。通过对辅助驾驶模型运行环境进行独立隔离，防止联网管理进程对智驾功能安全的干扰。



PART 3

华为乾崮智能汽车解决方案 网络安全技术

3.1 基础安全

车上各电子控制单元统称为ECU，典型的ECU是以微控制器（MCU）为主控芯片构建。智能汽车开始使用系统级芯片（SOC）构建ECU，包括车机、辅助驾驶系统等。乾崮智驾、鸿蒙座舱、乾崮车控、乾崮车载光均属于车载ECU之一。乾崮网络安全解决方案对ECU及其软件系统，实施一系列高质量的基础安全功能，贯穿ECU的规划、设计、开发、测试、生产、运行、维护全生命周期。

3.1.1 系统安全

1、安全启动

设备启动流程中的每一步，都包含对启动对象的数字签名校验，以确保设备在启动过程中加载并运行合法授权的软件。只有正确通过签名校验的镜像文件才可被加载并运行，包括启动引导程序、内核、基带、短距固件等镜像文件。在启动过程的任何阶段，如果签名校验失败，则启动流程会被终止。

2、安全升级

除了保证启动阶段系统软件的完整性及合法性，设备在OTA升级和近端刷写升级过程中，依然保证平台软件的完整性及合法性。系统软件更新时，会对升级包的签名进行校验，只有通过校验的升级包才被认为合法并安装。



3、可信执行环境

内含SOC的设备均支持可信执行环境技术。可信执行环境技术iTrustee基于TrustZone技术实现。TrustZone是硬件级别的安全，兼顾了性能、安全和成本的平衡。TrustZone技术将处理器的工作状态分为安全世界（简称TEE）和非安全世界(简称REE)。通过特殊指令在CPU的安全世界和非安全世界之间切换来提供硬件隔离。在安全世界，提供了对硬件资源的保护和隔离，包括内存、外设等，通过执行过程保护、密钥保密性、数据完整性和访问权限实现了端到端的安全，可防止来自非安全世界中的恶意软件攻击。

可信执行环境技术支持如下能力：

1）基础安全加固

- 可信执行环境全生命周期确保合法性及完整性，包含：启动、升级。
- 对镜像文件进行逆向分析是攻击者对目标发起攻击的重要手段，可信执行环境支持镜像防逆向保护。防逆向保护技术主要为镜像加密和符号表混淆。
- 可信执行环境支持防渗透，包括安全编译、地址随机化、栈保护、数据不可执行、代码段及函数指针只读。

2）安全服务

- 可信存储服务，提供关键信息的存储能力，保证数据的机密性、完整性。可信存储支持设备绑定，支持不同安全应用之间的隔离，可信应用仅能访问自己的存储内容，无法打开、删除或篡改其它应用的存储内容。HarmonyOS可信存储分为两种：安全文件系统存储与RPMB（Replay Protected Memory Block，重放保护存储区）存储，前者将密文存储到特定的安全存储分区，后者存储到eMMC（嵌入式多媒体卡）特定的存储区域，RPMB支持防删除、防回滚。
- 可信执行环境加解密服务支持多种对称、非对称加解密算法以及密钥派生算法，支持同一芯片平台相同密钥的派生，支持设备唯一密钥，支持标准的加密算法，为第三方开发存储和使用密钥的业务可信应用提供支持，并遵从Global Platform TEE标准。为提高安全性，可信执行环境内部的密钥生成和计算，均由独立的硬件芯片完成。



4、强制访问控制

支持强制访问控制特性，强制访问控制策略在设备启动时加载到内核中，无法被动态更改。该特性对所有进程访问目录、文件、设备节点等操作资源实施强制访问控制，对具有root权限的本地进程实施基于权能的强制访问控制，阻止恶意进程读、写受保护数据或者攻击其他进程，把被恶意篡改的进程对系统的影响限制在一个局部范围内，支撑上层应用实现各种安全防护。同时也支持Seccomp（Secure Computing Mode）特性，基于只读文件系统中的规则文件，对进程能够调用的系统调用进行限制，避免恶意应用通过使用敏感的系统调用对系统造成危害。

3.1.2 数据安全保护

数据生命周期包括**收集、存储、使用、传输、删除**这几个阶段：

收集：应用软件通过采集、直接生成、从其它设备接收或其它方式转入等方式产生数据的过程。

存储：数据在设备上存留的过程。

使用：数据在设备上被访问、处理等操作的过程。

传输：数据离开源设备、转移到目的设备的过程。

删除：数据在设备上被销毁，保证其不可被检索、访问的状态。

华为乾崮网络安全解决方案对数据进行全生命周期的安全防护，保证数据的机密性、完整性、可用性。

1、数据生命周期安全保护

- 1）密钥在TEE侧生成/存储/使用以及销毁，确保密钥明文不在REE侧存在。
- 2）数据存储支持全盘加密和文件加密，可以防止存储时窃取数据隐私。对于敏感个人信息，确保每次存储均进行加密，加密密钥安全存储在TEE中。





3) 数据使用阶段，使用自主访问控制（本地文件系统沙盒）以及强制访问控制能力，确保只有正确的应用才能够访问对应的数据。通过对文件加密密钥的管控，提供了文件保护增强能力，用于高风险等级数据在使用过程中的访问控制增强处理，确保有访问权限的软件才能够使用这些高风险数据。

4) 敏感数据出车时，在数据脱敏的基础上，进行数据加密传输。

5) 普通的恢复出厂设置操作，并不保证彻底删除保存在物理存储上的数据，为了提高效率，往往通过删除逻辑地址的方式实现，导致实际存储的物理地址空间没有清除，可以被恢复回来。

华为乾崮网络安全解决方案的恢复出厂设置，支持对存储数据的安全擦除。通过给物理存储器发送命令，进行覆写操作，完成底层数据擦除。擦除后数据是全0或者全1，确保用户的敏感数据不能通过软硬件手段恢复，能够保护用户设备转售、废弃后的数据安全。

2、多账号数据隔离

为了满足车辆共享场景，如共享给家人、临时共享给朋友、代驾、4S店保养维修等，系统构建了多用户多账号数据隔离机制，这种隔离机制是安全架构中的一个重要组成部分，旨在确保不同用户之间的数据和资源独立、安全。

在用户上车时，通过摄像头人脸识别，可以便捷地认证、切换用户，每个用户只能看到自己开车时形成的数据，进而保护账号内的隐私。

在将车辆交给代驾、洗车维修人员时，可以退出自己的账号，使用GUEST账号，从而实现将其他人临时用车时的数据与个人用车时的数据隔离，可防止其他人看到自己用车时留下的隐私数据。

3、敏感权限管理

华为乾崮鸿蒙座舱提供透明可控机制，帮助用户对敏感资源的访问行为可知可控，这些机制贯穿于应用整个运行期间，包括敏感数据或者能力被访问前、被访问中和被访问后各个阶段。

1) 隐私权限授权（访问前）：应用访问敏感数据或者能力前，需要申请相应的权限，对汽车关键的3个敏感权限进行了扩展设计和管理，包括麦克风、车内摄像头、位置信息权限；支持查看、撤销已授予应用的敏感权限。



2) 隐私指示器（访问中）：敏感数据或能力（例如麦克风）被应用持续访问时，通过状态栏显示实时提醒用户，便于用户感知应用访问行为；同时，车内摄像头支持物理遮挡，可以手动关闭，让用户更加放心，遮挡的方式可以是如下选择：

- ✓ 摄像头支持手动或自动伸缩/升降：摄像头伸出/下降时，摄像头露出，座舱应用可以获取摄像头数据。
- ✓ 通过拨片或滑片遮挡。

3) 权限使用记录（访问后）：鸿蒙智能座舱系统权限管理，包括位置信息、车内摄像头、麦克风，支持通过智能座舱查看敏感权限的访问记录。

3.1.3 可信互联通信

在网络中传输数据，包括智能车与云的连接，以及手机手表终端与云的连接，都采用安全传输通道来保证数据的安全，确保车、手机和服务端的网络连接中信息不会被窃取和篡改。安全通信采用国际标准或业界公认的安全协议，如TLS v1.2、TLS v1.3。在智能车中预置云的CA根证书，并预置由PKI签发的设备证书。车端强校验云端证书通过，且云服务器强校验车端的设备证书通过后，才能建立连接。

车云PKI使用安全性达到业界领先的根证书硬件加密机，根CA私钥保存在硬件加密机中，证书签发活动也在硬件加密机中完成，确保签名信息不可被假冒。端侧证书私钥保存在部件的TEE中，保证设备证书不可被仿冒。车内部件间的安全通信，也基于上述设备证书进行双向认证。诊断仪与车内网关，同样基于设备证书进行双向认证，以此完成接入认证，确保仅限合法的诊断仪可以接入车内进行诊断。





3.2 安全解决方案

3.2.1 远控端到端保护技术

远距离批量攻击，是对汽车危害影响最大的攻击之一。远程控制业务通过蜂窝网络，提供距离不受限制的车辆控制功能，典型的功能包括：门窗控制、空调控制、座椅控制、启动车辆、闪光灯鸣笛、充电控制等。如果远程控制业务被攻击者攻破，可能引起批量的远程攻击，攻击影响严重。

业界常见的做法，是通过TLS（传输层安全协议）安全通信进行保护。终端通过TLS通道将指令发送到车云，车云对指令进行必要的处理后，再通过TLS通道将指令发送到车端，车端执行远程控制指令，并反馈执行结果。

上述做法依赖TLS安全通道，存在两方面的潜在风险：

1、如果车云的API存在漏洞，被攻击者利用，可以通过车云仿冒控制指令，远程批量控制车辆，此类攻击危害很大。依据Upstream全球汽车网络安全报告的分析，超过70%的攻击是通过云远程实施。业界已经出现通过云远程批量控制汽车的攻击案例。

2、攻击者诱导用户终端信任并连接攻击者的设备，攻击者的设备作为中间人，篡改用户的远控指令，再发送到车云，车云和车辆会执行被篡改后的指令，引起非法控车。

华为乾崮网络安全解决方案的远控端到端技术，创新性地在终端（手机手表）与车辆之间做指令的签名和验签，在车主移动终端和车辆之间，形成密码学上的一对一绑定，签名密钥仅终端拥有，即使是云端也无法重新签名进行篡改。



1、完整性保证：端到端全程不可篡改，有效防范批量控车。只有车主及授权用户的设备可以发送远程控车指令，攻击者的仿冒设备不能远程控车。

2、通信通道中没有密钥传输，不会泄漏密钥。

3、用户无感完成指令签名，无需额外操作，保证安全的同时体验便捷。

3.2.2 一车一授权防刷机技术

智能部件必须配合合法的软件，才能保证功能正常可靠。如果智能部件的系统软件被改刷，攻击者可以在部件中植入控制软件、修改参数、绕过安全检查，非授权控制部件进而控制汽车，窃取敏感信息，对车辆安全威胁很大。

改刷系统进行攻击一直是一个难以彻底解决的问题，除了通过漏洞进行攻击以外，还包括使用具有车厂签名的版本进行攻击：

1、版本回退攻击

将一个旧版本刷写到商用车上，并利用该版本尚未修复的漏洞进行攻击，达成控制系统或者控制车辆。随着时间的推移，系统的漏洞总会被业界挖掘发现，正确的处理方式是尽快推出已修复漏洞的新版本，并推送OTA升级。但是攻击者可以特意刷入旧版本，实现回退攻击。由于旧版本也持有官方签名，一般的安全升级、安全启动方案不能防御回退攻击。





2、特权版本攻击

将调试版本、维修版本、测试版本刷写到商用车上，获得特殊权限，进而实现控制系统、控制车辆，或者篡改系统、改装车辆。攻击者可能通过社会工程学的方法，设法获得此类具有特殊权限的软件版本。车辆的研发、维修过程，包括部件底层软件开发、一级供应商集成开发、整车联合调试、众多售后服务中心，调试和维修版本在较多的环节和人员持有，有可能被攻击者获得。

华为乾崮网络安全解决方案的一车一授权防刷机技术，防止被篡改、替换、植入非法软件，包括防止软件版本回退到有漏洞的旧版本，进而防止使用已知漏洞进行攻击。

一车一授权技术方案，是在安全刷写、安全启动的基础上，结合芯片定制能力，实施的防刷机增强方案。该方案做到对每辆车、每颗芯片进行单独的版本授权，每次升级、刷写均需PKI签名授权，防止因为软件被非法刷改导致行车安全问题。

3.2.3 高安全虚拟机技术

智能座舱的功能越来越丰富，提供智慧语音、地图导航、音乐视频播放等功能，并且可以下载安装各类APP应用，引入了攻击风险。与此同时，在同一个座舱芯片和系统上，集成了仪表盘、AR-HUD、360度环视、电子后视镜等与车辆控制相关的功能。此类车辆控制相关功能，需要满足车规安全要求，保证可用性、实时性。如何保证座舱系统上车控相关功能的安全，免受APP漏洞、WIFI/蓝牙/互联网连接的攻击威胁，成为需要解决的问题。

华为乾崮网络安全解决方案使用高安全虚拟机，将座舱隔离为娱乐域和仪表域两个域。该虚拟机基于通过CC EAL 6+安全认证的微内核构建，并实施内核实时完整性保护技术。座舱将三方APP、人机交互、WIFI/蓝牙、网络连接等功能和外部接口部署在娱乐域，将车控相关功能部署在仪表域，两域互相隔离，以此保证车控相关功能的安全。



3.2.4 信任环车内通信安全技术

基于零信任的设计理念，车辆架构的内层部件，对外层部件的访问，进行接入认证和访问控制，防止外层对内层的渗透攻击。攻击者可能使用仿冒的零部件，替换正品零部件，或者使用旧件、报废件冒充新件安装到车上，达成攻击者的利益。此类攻击危及车辆功能安全，可能引起行车故障事故，扰乱配件管理，侵害消费者权益。防御的场景包括：

1、防仿冒件装到车上

来源：非授权的附厂件、三无水货件、私自加装的部件；

危害：危及车辆功能安全，扰乱配件管理，引起非法控制；

共同特点：仿冒件没有官方证书。

2、防非授权拆旧件装到车上

来源：报废车、偷盗车、修理车（如泡水车、火烧车）拆下的旧件，跨区跨地域串货件；

危害：欺骗消费者、引起故障、扰乱配件管理；

共同特点：拆旧件或者串货件具有官方证书。

3、窃取证书仿冒设备

来源：漏洞攻击、管理不善导致其中一个证书私钥泄露；

危害：可以批量仿冒设备，可以攻击所有车辆；

共同特点：窃取一个官方证书，或者虚构仿冒一个证书；

华为乾崮网络安全解决方案使用信任环技术防御上述攻击场景。

信任环是车内TLS通信的增强技术方案，在车内以太网连接部件之间建立接入认证和加密安全通信通道。信任环技术方案由数字证书保证，受信任环保护的智能部件均为原厂正品。





3.2.5 诊断安全技术

诊断接口（OBD口），是汽车用于检测维修的接口。通过诊断接口，可以读取故障码、修改车辆设置参数、刷写升级汽车的软件系统。汽车的大多数检测维修操作需要使用诊断接口。由于诊断接口功能强大，带来潜在的网络安全风险，例如非法修改汽车参数可导致功能异常、非法刷写软件系统可能植入控车后门、非法更换零部件带来行车安全威胁。利用诊断接口的攻击，需要进行至少一次的物理接触，如果达成软件植入，可以转化为远距离攻击。

诊断接口的业务操作及其风险：

场景	业务操作	风险分析	风险级别
研发制造	• 读、写设置参数	• 产线为封闭环境，整体风险可控	低
	• 诊断密钥分发	• 密钥传递过程存在泄露风险	高
维修	• 读取诊断所需数据	• 读取配置参数等信息	中
	• 写入修改配置和软件	• 如诊断仪丢失盗用，可刷写固件，修改配置	高
	• 获得诊断密钥接入车辆	• 诊断仪存在泄露密钥风险 • 所有车辆密钥相同，形成规模化攻击	高

由于售后维修中心众多，维修人员众多，诊断仪容易出现多人共用、出借、丢失的情况，攻击者也可能通过社会工程学获得诊断仪，进而可以对车辆进行诊断、修改。

在以往的诊断方案中，诊断安全是依赖诊断密钥的机密性，持有诊断密钥的设备就可以对车进行诊断修改。该诊断密钥是对称密钥，任何一个地方泄露，将导致本车型所有车辆可以被诊断修改。由于业务开展的需要，该诊断密钥需要保存在诊断中，也保存在部件中，两种设备都可能被攻击者通过逆向分析的方式获得密钥。

华为乾崮网络安全解决方案包含了诊断安全技术方案，为诊断接口提供强有力的安全防御，可以保证只有官方授权店可以维修和更换智能部件，防止攻击者非法篡改汽车参数。实施的有效安全措施包括：



1、使用密码学公认的标准安全密码算法

摒弃了业界过往常见的私有简易算法，华为乾崮网络安全解决方案从一开始就坚持全部密码算法使用公认标准安全算法，可以从理论上证明加密算法不会被攻破。同时，对密钥的生成、分发、保存、运行、销毁，实施全生命周期的安全方案，确保密钥不泄露。

2、关键部件首创实现一车一件一密钥，一车泄露密钥不会影响其他车辆

业界常见的做法是一部件型号一诊断密钥，即不同型号的部件密钥不同，同型号部件的诊断密钥相同。由车企将密钥分发给部件供应商，由部件供应商安全预置到部件中，随部件发货到整车产线。此时，车厂和部件供应商均持有该密钥，任何一方以及分发过程中保管不善，泄露密钥，所有使用该型号部件的车辆都可能被攻击。另外，攻击者对市面上其中一辆车进行逆向分析，一旦成功获取密钥，可以对市面上整批车辆实施攻击。

华为乾崮网络安全解决方案的诊断安全技术，在业界首次做到每辆车每个部件诊断密钥不同，一辆车中的密钥即使被攻击者获得，不会对市面上其他车辆造成安全风险。并且，诊断密钥在车内TEE安全保存，在维修时无需将密钥下发到诊断仪，进一步减少了密钥泄露风险。

- 1）每个智能部件诊断密钥不同，减少泄露的影响。
- 2）维修时，诊断仪不需要获得诊断密钥，进一步降低泄露风险。
- 3）诊断密钥在车内由TEE或者HSM（硬件安全模块）安全存储，防止逆向、物理攻击窃取密钥。
- 4）密钥传递全程安全加密，由IT设备自动化对接传递，不经过人工传递。

3、每次接入诊断接口进行诊断操作，均需授权

只有合法的诊断仪可以检修智能部件，可以防止未经授权的无资质人员篡改汽车。每次需要检修智能部件，需要在线认证诊断仪设备，并且指定只能在指定的时间内，对指定的车辆开展检修工作。



PART 4

华为乾崮智能汽车解决方案 网络安全法规标准遵从、测评与认证

引望始终将技术标准制定视为推动行业进步的关键，以开放协作驱动技术落地，以自主创新赋能产业升级。凭借深厚的技术积累，通过研发投入、专家协同等方式推动标准完善，确保其科学性与实用性。华为曾共同牵头制定了汽车行业首个信息安全推荐性国家标准GB/T 40861-2021《汽车信息安全通用技术要求》，并于2021年联合中国汽研共同提出“管理体系+产品测试”为理念的汽车信息安全能力评价方法，发布了汽车行业首个《智能网联汽车信息安全测试评价白皮书》，也为后续汽车领域权威测评IVISTA的专项测试提供了指导依据。

引望将标准遵从融入企业社会责任。无论是中国还是欧盟，均建立了严苛的合规体系。标准是产业发展的“灯塔”，引望的投入不仅为自身构筑技术护城河，更以开放姿态为汽车智能化进程提供企业智慧。未来，引望将继续以标准为纽带，与各界共创安全、共赢的智能网联汽车产业。

4.1 GB 44495遵从

由工业和信息化部归口，中国汽车技术研究中心有限公司（以下简称“中汽中心”）等单位起草的GB 44495-2024《汽车整车信息安全技术要求》国家强制性标准已于2024年8月23日正式发布，并将于2026年1月1日实施。标准规定了汽车信息安全保障要求、信息安全一般要求、信息安全技术要求、检验与试验方法以及同一型式判定。该标准是国内首个将信息安全保障要求作为强制要求的标准，明确规定企业需制定详尽的流程和规范，以确保风险管控工作有序开展，有效减少人为错误和疏漏。尽管保障要求的建立周期较长，但它对于规范和加强智能网联汽车的信息安全管理，以及为企业在汽车整车信息安全建设方面提供了极大的帮助。

华为是GB 44495-2024《汽车整车信息安全技术要求》标准的主要参编单位之一，凭借在通信技术、网络安全和数据加密等领域的优势，为标准的科学性和实用性提供了技术支撑，与中汽中心等机构共同完成了外部连接安全、通信安全、软件升级安全及数据安全等标准要求的设计。



4.2 UN R155、ISO/SAE 21434遵从

WP.29是联合国欧洲经济委员会（UNECE）内陆运输委员会（ITC）的附属机构，目标是在全球范围内协调或制定适用于汽车领域的技术法规，并通过发起和采取相关法规行动，加强全球范围的经济关系，其系列法规涉及到轮式车辆，其子系统和零件的安全和环保性能。

2020年6月25日WP.29颁布了第155号法规（以下简称“UN R155”），该法规是国际汽车网络安全领域首个具有约束力的准则，标志着汽车网络安全正式进入了强制监管范围。

UN R155中管理体系的部分聚焦组织管理、风险管理、供应商管理、漏洞及事件管理等方面，对车辆制造商（OEM）在网络安全管理体系的建设作出了相应的规范要求。该管理体系的规范旨在要求车辆全生命周期中包括产品的开发、生产、销售运维等各个阶段均有流程体系指导，并能支持OEM有效发现和控制风险。UN R155的强制实施，一方面为OEM在网络安全技术研发方面确立了更清晰的法规指引，以保证终端用户的使用安全；另一方面也对有计划出口至《1998年协议书》相关协定国家和地区的OEM而言提出了更加严峻的准入挑战。

车辆网络安全型式认证（VTA）则是针对车型在开发过程中具体工作进行网络安全审查，旨在保证实施于车辆的网络安全防护技术在进行审查认证时足够完备。

ISO/SAE 21434是由SAE和ISO联合起草发布的标准，是UN R155的关键支撑标准。该标准于2021年8月31日正式发布，定义了汽车电子电气系统的网络安全风险管理要求，覆盖概念、开发、生产、运维、报废等全生命周期各个阶段。符合ISO/SAE 21434标准可以帮助汽车制造商和零部件供应商满足联合国汽车网络安全法规的管理体系要求。

2021年，华为智能汽车解决方案BU正式获得汽车网络安全ISO/SAE 21434符合性证书，成为全球首个通过德凯ISO/SAE 21434认证的智能汽车解决方案一级供应商，有效支撑整车企业通过R155认证。



ISO/SAE 21434标准认证的通过，表明华为智能汽车解决方案BU在产品研发、采购、生产、运维等全生命周期流程完全满足标准的要求，具备为汽车行业客户提供符合业界最佳实践的网络安全能力。

4.3 C-ICAP测评

C-ICAP全称是中国智能网联汽车技术规程（China Intelligent Car Assessment Program），是中汽中心汽车测评管理中心联合行业经历三次修订，于2024年6月发布的独立、公正、专业的智能网联汽车整车性能综合评价，为提升智能网联汽车用户满意度，引领智能网联汽车企业创新能力，培育智能网联汽车品牌影响力提供专业平台。隐私保护作为其中一个重要的单元，与辅助驾驶、智慧座舱共同构成了C-ICAP（2024版）的全部测评项目。其中，隐私保护单元测评项目包括网联通信守护和个人信息守护，前者旨在对智能网联汽车防入侵、防篡改能力进行考查，后者旨在对智能网联汽车防偷窥、防窃取能力进行考查。

► **网联通信守护**：分为基础安全守护和攻击防御守护，前者从消费者更容易感知的角度（车载蓝牙、车载Wi-Fi、NFC钥匙、射频遥控钥匙、手机APP远程解锁）进行包括漏洞检测、模糊测试在内的常规网络安全检测；后者则偏向真实的黑客攻击，使用中汽中心自行构造的通用攻击库、舆情用例库、黑客攻击用例库以及组合节点的攻击方式对整车开展进阶测试。

► **个人信息守护**：针对车端敏感个人信息（包含车内视频图像、车内音频、车辆位置等）、车端一般个人信息（包含手机号、账号、车牌号等）从用户权益保护角度开展告知情况、授权提醒、数据删除的测试；此外，手机端个人信息（聚焦控车APP）也在本技术规程测评范围内，对于APP是否存在信息过度收集/未授权收集个人信息、是否存在未授权插件嵌入、是否存在未授权向其他第三方共享个人信息的情况进行逐一测试。





C-ICAP（2024版）隐私保护测评规程根据各项得分与权重计算出得分率从而获得最终的评价，当得分率超过92%（含）则为最高的“五星”评价。

在C-ICAP（2024版）隐私保护测评规程制定过程中，考虑到消费者可能存在的手机更换、开通授权账号等日常使用场景，华为提出对控车APP增加二次验证的要求，进一步降低车辆被非法控制的风险。

在C-ICAP（2024版）正式发布后，鸿蒙智行享界S9也及时开展了摸底测试与正式测试，最终得分率达到97.4%，获得隐私保护板块最高级（五星）评价。未来，引望仍将深度参与C-ICAP隐私保护测评规程后续版本的讨论与制定工作，践行消费者个人权益保护，把隐私安全带入每一辆车。

4.4 IVISTA测评

IVISTA全称是中国智能汽车指数（Intelligent Vehicle Integrated Systems Test Area），是中国汽车工程研究院股份有限公司（简称“中国汽研”）自2017年打造的全球首个面向消费者的公平、公正、专业、权威的智能网联汽车第三方测试评价体系。2024年7月，IVISTA正式发布了网联智能与隐私安全专项测评（2023版），从网络安全、隐私安全及网联智能功能三方面对智能网联汽车面临的日益严峻的网络安全和数据安全问题开展量化评价，牵引行业提升智能网联汽车网络安全水平，保障产业健康有序发展。

► **网络安全：**对数字钥匙的防重放、防中继功能，导航定位的防欺骗、防干扰功能，远程控车的APP安全、通信安全功能，车端接口（USB、OBD、充电接口等）的防病毒、防提权功能，Wi-Fi/蓝牙的防破解、防钓鱼功能进行了全方位的严格测试，此外，车辆还可根据自身情况选择参与极限攻防测试，接受专业团队实战化的攻击。



► **隐私安全：**对座舱中的传感器（摄像头、麦克风等）的默认不开启、采集数据默认不出车功能，对处理敏感个人信息的显著告知、单独同意、自主设定同意期限功能，对个人行权渠道的合理有效性，对数据的防境外直接传输进行了全面的验证，此外测评鼓励车辆配备更多增强座舱隐私的安全措施，并对满足条件的措施进行额外加分。

► **网联智能：**对数字钥匙功能的丰富度，远程控车功能的丰富度和成功率，哨兵模式的有效性和能耗进行评价，鼓励车辆具备更多智能化人性化的功能。

IVISTA网联智能与隐私安全专项测评（2023版）采用各项得分累加制衡量最终的评价，当网络安全和隐私安全板块得分分别超过38分（含）及11分（含）时，则可获得“G”评价（最高等级），当且仅当网络安全和隐私安全都获得“G”评价时，整体才能达到“G”评价。

在IVISTA网联智能与隐私安全专项测评（2023版）规程的制定过程中，华为提出保护个人隐私的账号隔离、主驾隐私模式、隐私声盾等安全措施，实现防跟踪、防监视、防泄漏的驾驶环境，贡献企业智慧。

在IVISTA网联智能与隐私安全专项测评（2023版）正式发布后，鸿蒙智行智界S7 2024款长续航Max版本以总分第一的成绩（网络安全得分57，隐私安全得分21）在第一批测评中脱颖而出斩获最高“G”评价，成为智能汽车网络与隐私安全防护的标杆车型。未来，引望将持续参与IVISTA网联智能与隐私安全专项测评后续的修订及新版本制定工作，为行业创造更多价值。





4.5 CC EAL安全认证

CC（Common Criteria）认证是一种信息安全认证，是全球IT行业中最具影响力的认证之一。CC认证依据信息技术安全评估通用标准对IT产品的安全功能和安全保障能力进行全方位评估，涉及产品的设计开发、安全功能、交付管理等方面。目前有31个国家签署了CCRA（Common Criteria Recognition Arrangement）互认协议，有17个国家可以颁发证书，31个国家采纳和认可，是全球广泛认可的权威安全认证。

EAL（Evaluation Assurance Level）指评估保障等级，是IT产品或系统在CC安全评估下的数字等级。CC预定义了7个保障级，从EAL1到EAL7，每一级均需评估7个功能类，分别是配置管理、分发和操作、开发过程、指导文献、生命期的技术支持、测试和脆弱性评估。满足了保障级相应的安全保障要求，就达到了相应的保障级。等级越高，表示通过认证需要满足的安全保证要求越多，系统的安全特性越可靠。

CC认证及其评估保证级别EAL，主要用于评估产品或解决方案的安全性、可靠性以及隐私的保护能力，是目前公认的全球认可度和权威性最高的产品安全认证，可作为企业信息技术解决方案建设时产品安全性评估的重要依据。

华为VOS（Vehicle Control OS，智能车控OS）安全车控操作系统于2022年正式通过了全球独立安全机构欧洲SGS Brightsight 实验室的安全评估，并通过了CC评价体系中最严格的AVA.VAN.5的脆弱性评估要求，获得全球权威信息技术安全性评估标准CC EAL4+AVA.VAN.5级认证证书，成为国际上首个通过CC EAL4+级别认证的车控操作系统。

随着汽车电动化、网联化、智能化和共享化的发展趋势，智能网联汽车为用户提供了卓越的使用体验，但同时也面临着日益复杂的网络安全挑战。智能网联汽车的安全性高度依赖于其电子控制单元（ECU）的安全性。一辆现代化的智能网联汽车通常配备数百个ECU，这些ECU负责控制车辆的行驶状态及实现多种智能化功能。因此，整车的网络安全状况与每一个联网ECU的安全性密切相关。



华为自主研发的VOS是一套兼容AUTOSAR CP标准的基础软件平台，专为车辆的各种ECU产品设计。华为VOS为应用程序开发人员提供了一个软件平台，用于开发控制ECU的应用程序，并通过对象识别、对象管理、应用管理、服务保护、存储保护等安全机制，为基础软件平台上的可信应用提供全面的安全保障。作为安全车控操作系统，华为VOS已经满足ECU的高等级安全要求，确保系统整体达到相应的安全防护能力。

4.6 车联网云平台CCRC认证

中国网络安全审查认证和市场监管大数据中心（China Cybersecurity Review, Certification and Market Regulation Big Data Center, CCRC）是国家市场监督管理总局直属正司局级事业单位，依据国家《网络安全法》及国家有关强制性产品认证法律法规，负责实施网络安全审查和认证。

CCRC开展的IT产品信息安全认证业务，是依据信息技术安全评估准则和相关技术要求，对IT产品的安全性进行评价，旨在保护用户信息安全，维护用户利益。生产企业的IT产品获得信息安全认证证书，表明该产品符合相应的标准和技术要求。同时，认证过程也将有效促进服务提供方完善自身管理体系，提高服务质量和水平，引导行业健康规范发展。

2024年，华为联合CCRC及检测机构，首次制定和发布行业标准CCRC-TR-147-2024《车联网云服务系统安全技术要求》，填补了国内车联网云服务系统权威安全认证标准的空白。同年，华为乾崮智能车云服务取得业界首个车联网安全和数据安全最高等级认证，符合标准GB/T 18336-2015《信息技术 安全技术 信息技术安全评估准则》、CCRC-TR-147-2024《车联网云服务平台安全技术要求》（三级）、产品认证实施规则（CCRC-IR-001:2020）的要求。其中，车联网云服务平台安全的第三级要求包括安全运营要求、隐私中心要求、数字钥匙要求、远程诊断要求、远程监控要求、远程控制要求、升级服务要求、通用安全要求。未来，引望将持续推动车联网云服务的网络安全标准建设，积极参与制定和完善相关安全标准，推动技术创新与最佳实践的分享，以应对不断变化的网络安全挑战，构建一个更加安全、可靠的车联网生态系统，为用户提供更高水平的安全保障，促进智能网联汽车的可持续发展。



4.7 汽车隐私保护测评规范

智能网联汽车中摄像头、麦克风等传感器的普及，给个人隐私保护带来了巨大的挑战。为了增强消费者对汽车隐私保护能力的信心，中国汽车工业协会（以下简称“中汽协”）与中国网络安全产业联盟汽车网络安全工作委员会（以下简称“车安委”）联合国家级技术机构制定了一套汽车隐私保护能力测评规范。汽车隐私保护能力测评针对车辆所需使用敏感个人信息的相关功能，围绕隐私声明、车内外摄像头、麦克风、定位信息和用户体验计划展开，对包括车外人脸信息等匿名化处理、座舱数据不出车、座舱数据默认不收集以及处理个人信息显著告知4项要求进行全面评估，并将测评结果借助中汽协区块链进行存证，确保测评结果的真实可信。

华为前期已支撑阿维塔12车型通过汽车隐私保护能力测评，并获得了《汽车隐私保护标识证书》及“隐私保护标识”使用授权，提升了用户对隐私保护能力的感知度。目前，行业仅11款车型获得了上述证书及标识使用授权，彰显出其出众的隐私保护能力，为行业树立标杆。

未来，引望将继续在中汽协及车安委的带领下，支撑更多合作车企通过汽车隐私保护能力测评，为提升行业隐私保护水平，维护国家数据安全贡献力量。

PART 5

华为乾崑智能汽车解决方案 全生命周期安全保障

从汽车的概念设计阶段开始，一直到车辆报废，在整个生命周期的各个阶段都要充分考虑网络安全问题。这意味着在每个环节，包括需求规划、设计、实施、集成、验证、确认、生产、运行、维护以及最终的报废处理，都需要实施相应的网络安全措施，确保车辆在整个使用过程中都能抵御各类网络攻击，保障用户数据安全和车辆行驶安全。华为乾崑网络安全解决方案，覆盖了汽车全生命周期的网络安全保障。



图6：全生命周期安全保障

前文已经对安全规划、安全设计及技术方案进行了介绍。本章重点阐述安全工程（重点支撑安全设计、安全开发工作）、安全测试、漏洞管理、应急响应环节。

5.1 安全隐私可信工程

恪尽职守、夯实信任：引望将对网络和业务安全性保障的责任置于公司的商业利益之上，坚持投入，开放透明，全面提升软件工程能力与实践，构筑网络韧性，打造可信任的高质量产品，保障车联网和用户隐私的安全。

可信任不仅仅是产品外在表现的结果，更是产品内在实现的过程，是结果和过程双重可信的高质量。每一位具备可信价值观的员工，基于引望可信任过程相互协作创新，开发出具备可信特征的产品，给客户提供可信的高质量产品，并持续改进。

我们从系统工程的行业共识出发，基于可解释、可落地、可验证和有相当业界共识基础的四个原则定义了智能汽车可信框架。

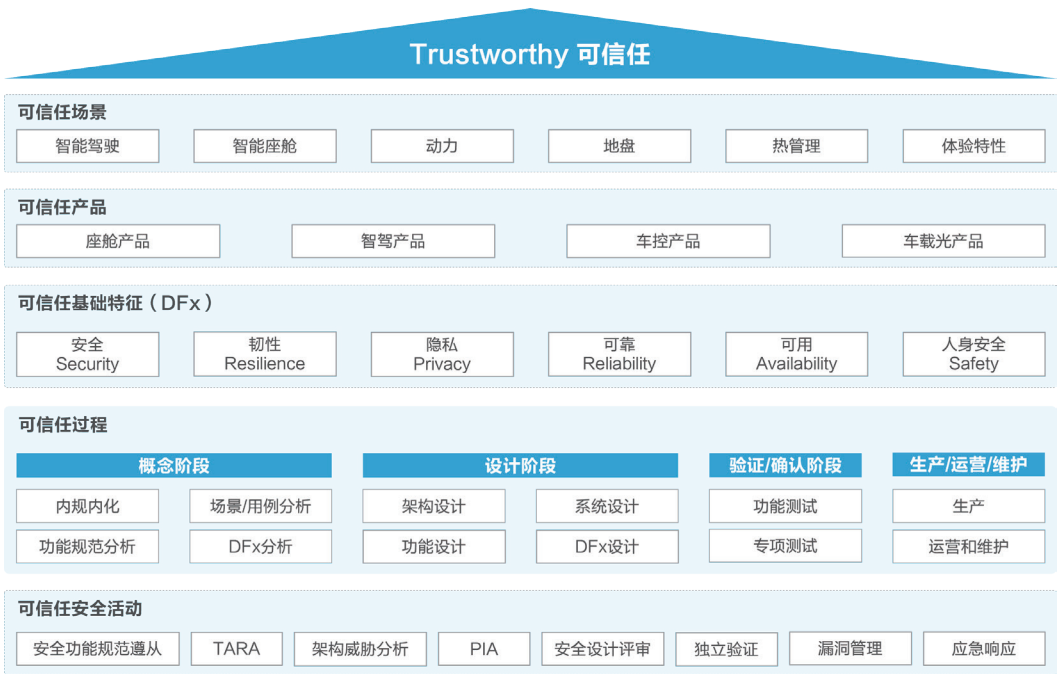


图7：智能汽车可信框架

我们要在每一个智能汽车产品和解决方案中，全面融入可信任安全活动，包括：

1、安全功能规范遵从：将外部标准法规和业界最佳实践进行解读和提炼，融入安全功能规范，并确保功能规范在产品中有效落地，以满足合规性及其他相关要求。

2、TARA分析（威胁分析与风险评估）：依据ISO/SAE 21434的要求，通过资产识别、损害场景识别、威胁场景识别、攻击路径分析、风险评估和风险处理等步骤，识别潜在威胁、分析其可能性和影响，并制定相应的防护措施。

3、架构威胁分析：对产品系统架构的安全威胁进行全面的分析和建模。从攻击者视角识别价值资产、暴露面和攻击路径，评估相关风险，并基于架构元素构建威胁模型，从而提升系统的网络韧性。

4、PIA分析（隐私风险分析）：对隐私风险进行管理的一系列活动的总称。通常以信息流（Information Flow）为工具，对个人可识别信息（PII）在业务处理活动中可能存在的隐私风险进行识别、分析、评估和处置。

5、安全设计评审：通过自检或者同行评审的方式，对设计方案进行安全性检查。包括安全功能规范检查、安全红线检查、威胁分析过程检查等，识别设计中存在的安全问题。

6、独立验证：参考5.2“VCSL网络安全实验室”章节。

7、漏洞管理：参考5.3“漏洞管理”章节。

8、应急响应：参考5.4“应急响应”章节。



5.2 VCSL网络安全实验室

VCSL（汽车网络安全实验室）是引望的安全测试实验室，面向引望产品解决方案，以及合作车型，按照ISO21434以及IPD-IAS流程，参考汽车行业安全隐私标准以及ICT领域网络安全和隐私保护的优秀实践，开展安全隐私测试，保障上市产品无重大安全隐私风险遗漏，保障合作伙伴以及终端用户的网络安全和隐私保护权利不受侵犯。

VCSL网络安全实验室的主要活动包括：

1、产品量产前质量保障：对引望交付的产品和解决方案，在量产前开展安全隐私测试，测试基线来源于国内外汽车行业强制标准、通过TARA洞察的安全需求，以及ICT领域安全隐私的优秀实践，通过FUZZ测试、渗透测试、验证测试和确认测试，达成安全目标，交付给客户高质量的产品。

2、独立验证：由独立于产品开发团队的“VCSL网络安全实验室”对安全设计和开发进行测试验证，确保设计方案和代码实现符合安全隐私要求。验证手段包括合规验证、渗透测试、Fuzzing测试、漏洞扫描等专项测试活动。

3、联合车企开展整车测试：对引望合作的车型，联合车企开展整车合规测试和渗透测试，满足国家强标要求以及海外准入要求，整车具备行业领先的安全隐私竞争力，满足IVISTA、C-ICAP等测评要求。

4、合作伙伴安全测试赋能：通过合作项目，对合作车企进行安全测试赋能，帮助车企引入ICT领域先进的安全隐私测试方法，保障产品上市质量和IT系统运营安全。

5、构建汽车行业安全影响力：参与国内外安全隐私测评、安全认证标准制定，和国内外测评机构共同打造全行业权威的汽车安全隐私测试标准；通过汽车安全峰会、学术界的合作，共同提升汽车行业安全测试能力。



引望VCSL致力于和全行业共同缔造汽车的安全测试能力不断提升，通过合作车型打造行业标杆；持续追踪ICT领域最新的安全技术发展，通过不断迭代，使能汽车行业安全水平持续领先。引望VCSL保持合作开放，积极和全行业优秀的安全实验室、学术界、测试机构等进行沟通和合作，共同参与汽车安全生态构建。

5.3 漏洞管理

5.3.1 漏洞管理是保障网络空间安全的重要工作

近年来，网络空间安全威胁不断升级，网络勒索、供应链安全等事件频发，造成严重损失，漏洞仍然是安全事件的主要触发点之一。围绕产品全生命周期构筑全供应链的端到端漏洞管理，是降低现网风险、保障业务连续运行的重要手段。随着全球产业数字化转型的深入，网络安全攻击向常态化、自动化发展，高风险漏洞和安全事件成为立法监管和相关技术孵化的加速剂，中英欧美等国陆续发布相关法规统筹漏洞管理，漏洞管理已成为国家网络安全战略的重要部分。漏洞管理的成熟度是企业数字化治理水平和软件工程能力的直接体现，与企业可持续发展密切相关，企业要和各利益方持续协同管理漏洞，否则可能面临系统瘫痪、信息泄露等风险，损害企业资产和声誉，甚至影响企业的长远发展。



5.3.2 漏洞管理是产业链的共同责任，需要协同合作来共同解决

漏洞管理涉及端到端供应链，如供应商（包括开源社区）、设备商、车企、最终消费者等多个利益相关方，也涉及产业全生命周期的漏洞感知、验证、修补、披露、现网漏洞风险消减等多个环节，如何确保漏洞信息在各利益相关方的及时、准确、安全的交互是行业挑战。同时企业内部产品、平台、组件等现代软件大规模的多组织协同开发，进一步加剧了漏洞管理的复杂性。漏洞管理需要供应链上下游协同，促进各责任方履责，以开放合作的态度，在供应链上下游构筑持续信赖的合作关系，不断提升信任与能力，共同消减漏洞带来的网络安全风险。

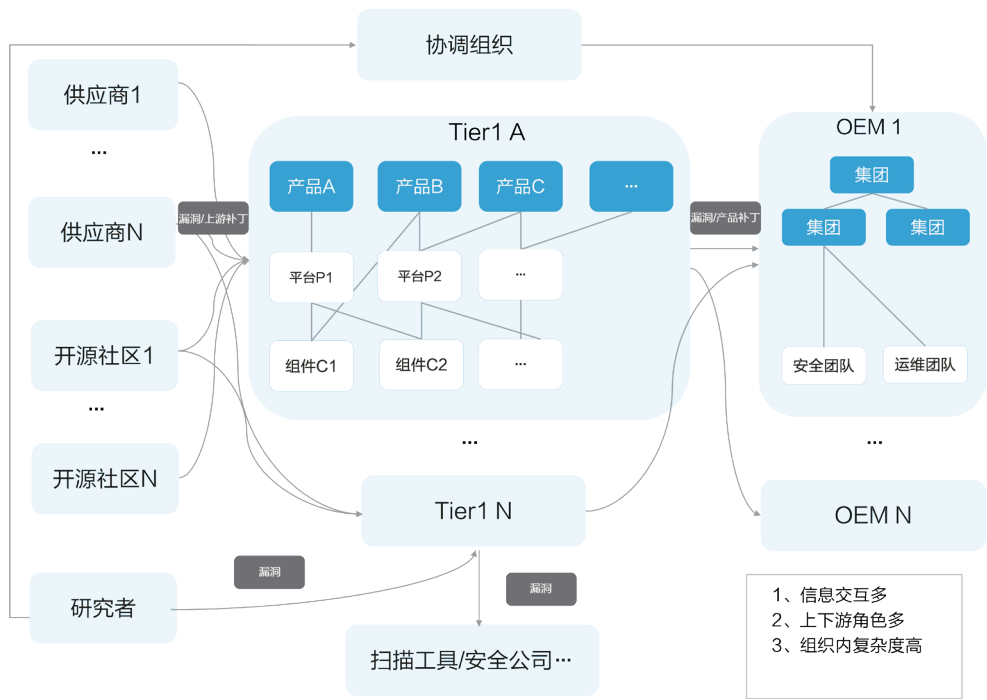


图8：漏洞管理的协同/协作关系图

5.3.3 引望遵循行业标准和最佳实践建立漏洞管理体系和工程能力

引望漏洞管理理念：引望将“**端到端的全球车联网安全保障体系**”作为公司的重要发展战略之一，从政策、组织、流程、技术和规范多方面建立可持续、可信赖的漏洞管理体系，与外部利益相关方开放协同，共同应对漏洞的挑战。引望对漏洞管理提出五项基本原则，确定七个漏洞处置阶段和相应的行为准则，指导业务部门开展漏洞管理活动。



图9：漏洞管理理念架构

5.3.3.1 漏洞管理目标

更好的支撑客户现网漏洞风险消减，引望将漏洞管理的目标具体分解成以下三个主要方面：

1、负责任披露：对购买引望产品和解决方案的客户，建立漏洞披露与沟通机制，支撑客户对漏洞风险决策。



2、减少和消减漏洞：建立全量、全生命周期、端到端漏洞管理机制，实现漏洞及时感知、排查、规避和修补，支持客户风险消减。

3、协同管理：明确与供应商和客户的漏洞管理协同机制，共同协作消减漏洞风险。

5.3.3.2 漏洞管理原则

1、伤害和降低风险：减少或消除漏洞给客户带来的伤害，降低引望产品、服务的潜在风险，既是漏洞管理的愿景，也是漏洞处置&漏洞披露时所遵循的价值指引。

2、减少和消减漏洞：漏洞不可避免是业界已经形成的共识，但我们一直在努力：

- 1）采取措施减少产品和服务中的漏洞；
- 2）一旦发现产品和服务中的漏洞，及时向客户/用户提供风险消减方案。

3、主动管理：漏洞问题需要供应链上下游协同合作来解决，我们会主动先厘清边界再履行责任，基于法规、合同和公开标准要求等，构建漏洞管理体系主动开展漏洞管理。

4、持续优化：网络安全是持续演进的动态过程，威胁和攻击在更新，防守也要持续更迭。我们将持续优化，不断借鉴行业标准和业界优秀实践，提升漏洞管理的成熟度。

5、开放协同：我们将秉持开放合作的态度，如定期召开安全沙龙，加强与供应链和外部安全生态的联接，并加强与利益相关方的协同，构筑可信赖的合作关系。



5.3.3.3 漏洞处置关键阶段

引望致力于提升产品和解决方案的安全性，全面支撑客户业务的安全运营。参考ISO/IEC 30111、ISO/IEC 29147标准建立完整的漏洞处置流程，持续为产品安全性保驾护航，并建立开放协同生态，共同应对漏洞风险挑战。

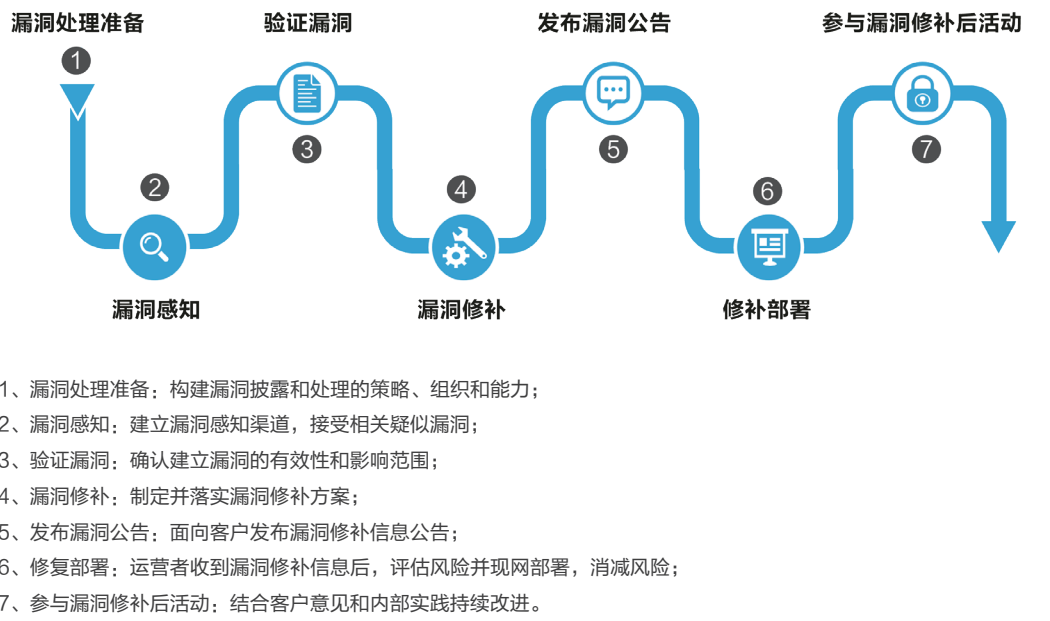


图10：漏洞处置的7个关键阶段

5.3.3.4 漏洞管理行为准则

为了更好地指导各业务部门开展漏洞管理活动，针对漏洞处置过程的关键活动，提出漏洞管理行为准则：

1、在漏洞处理准备阶段，需建立漏洞处理的组织与流程，明确漏洞处置目标，并与外部建立沟通渠道。





2、在漏洞感知阶段，应建立广泛的内外部漏洞感知渠道，全面感知与引望产品相关的漏洞信息，以便及时对漏洞进行验证和修补。

3、在验证漏洞阶段，验证感知到的漏洞对生命周期内版本的影响情况，在此期间保持与漏洞上报者持续、清晰的沟通。

4、在漏洞修补阶段，应积极采取行动对漏洞进行修补&消减。同时基于漏洞严重等级，对漏洞进行分级的响应，确保所有漏洞都得到合理处置。

5、在发布漏洞公告阶段，需遵循受影响告知的原则，即漏洞信息仅披露给受影响的利益相关方，同时做好内外部协同。

6、在修复部署阶段，在支撑客户或自身修复部署时，开展高风险漏洞的主动沟通，倡议客户/用户积极开展补丁管理建设，鼓励客户/用户主动关注安全更新公告，及时感知并消减。同时倡议监管机构建立定期审视机制，确保客户/用户及时更新最新的漏洞修补补丁。

7、在参与漏洞修补后活动阶段，产品团队应开展漏洞的根因分析来改进安全活动。另外通过外部沟通、产业链上下游的协同，持续发展和培养全员安全意识与能力。



5.3.4 引望漏洞管理框架与实践

为支撑客户现网漏洞风险消减的管理目标，引望建立漏洞管理框架，包括政策、流程、工程、文化与组织，并持续优化各领域能力，支撑漏洞管理的高效、有序开展。



图11： 引望漏洞管理框架

5.3.4.1 全量漏洞管理

资产管理是漏洞管理的基础，如何确保资产信息的完整性和准确性是行业挑战。引望产品涉及多个产业，拥有部件、云服务等诸多不同类型的产业资产，获取全面且更新的资产清单是个挑战，识别资产中准确的开源软件、三方件等成分也同样是业界难题。引望明确产业商业领袖作为漏洞管理的第一责任人。在产品立项时，须进行资产注册，从源头确保资产清单纳入管理；在开发过程中，通过软件工程能力（全量源码构建等）确保使用的资产来源已源头纳管，包含使用的平台、开源软件、三方件等。基于获得的全量软件资产，形成产品版本的全量漏洞视图，在此基础上开展验证、修补等漏洞管理相关的业务活动。



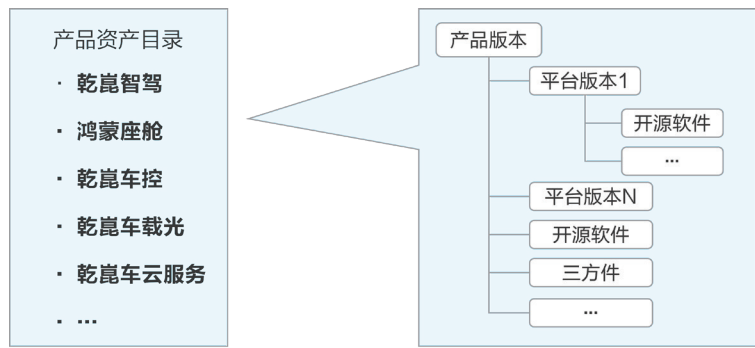


图12：全量管理示意图

5.3.4.2 全生命周期漏洞管理

引望将“在设计中构筑安全、在流程中构筑安全、在运作中构筑安全”的理念贯穿在产品的全生命周期，通过流程确保网络安全保障措施在各阶段有效落地，提升产品的安全竞争力。



图13：网络安全Build-in IPD流程

“集成产品开发流程”（IPD, Integrated Product Development），从概念、计划、开发、验证和发布阶段，将威胁建模、安全设计、安全开发、安全测试等Build-in流程中，从机制上减少漏洞的引入，在版本发布时，向客户交付漏洞风险最小化的产品版本。漏洞管理基于产品/软件版本的生命周期里程碑进行管理，引望对停止服务与支持（EOS）前所有产品版本的漏洞进行管理，根据不同生命周期阶段修补策略，在停止全面支持（EOFS）前发布漏洞修补方案（包括缓解措施、补丁/版本等），支持客户例行消减现网漏洞风险。引望漏洞管理中心持续识别公众广泛关注且已活跃被利用的“高风险”漏洞，并加快漏洞响应的过程，在24小时内发布SN（Security Notice，安全公告），通知受影响的客户。当有漏洞修补方案后，引望会通过SA（Security Advisory，安全通告）为相关受影响客户提供风险决策和消减支持。

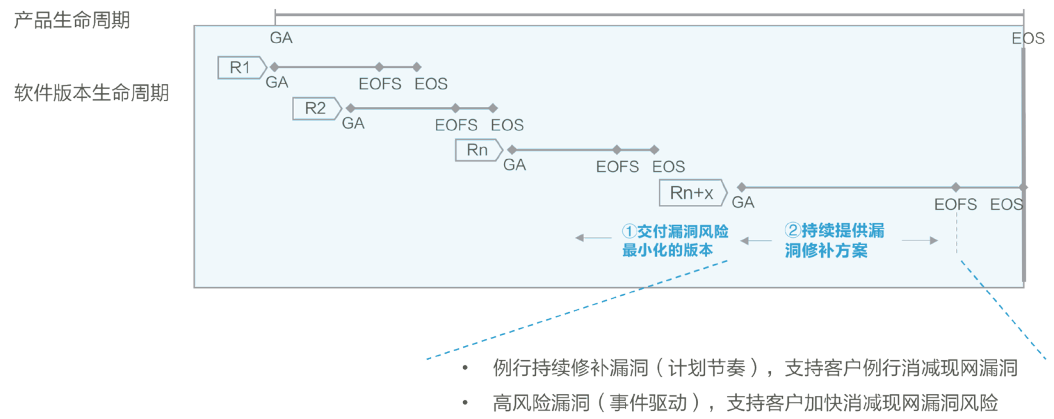


图14：全生命周期内漏洞管理示意图

5.3.4.3 全供应链管理

引望拥有复杂的产品、服务与解决方案组合，服务于全供应链中不同类型的客户，承担了包括供应商、设备商/集成商和车企等在内的多个角色，角色的多样使得引望有机会站从不同视角认识到供应链管理协同的重要性。

作为供应商，建立漏洞接收和感知渠道，以保证漏洞的及时、准确感知。同时，面向下游设备商/集成商，建立漏洞披露网站和主动沟通的渠道，通过持续提供漏洞修补方案，支持下游设备商/集成商对修补方案的集成和发布。



作为设备商/集成商，同供应商建立漏洞接收、披露、协同与响应机制。同时，开展漏洞奖励计划项目，鼓励安全研究者、组织等上报产品疑似漏洞。

面向客户，建立漏洞披露网站和主动沟通的渠道，以SA/SN、RN（Release Notes）的方式，向客户进行漏洞披露，支持客户知情决策和现网漏洞风险消减。

作为运营者，同设备商/服务提供者建立漏洞接收、披露、协同与响应机制。在现网资产管理基础上，持续开展漏洞感知、评估和现网修复等活动，将现网漏洞风险控制在可接受的水平。

5.3.4.4 漏洞管理平台

为了保障各个产业履责，支撑客户漏洞风险消减，引望漏洞管理中心作为漏洞信息和组织的汇聚点，建立集团统一的管理平台，通过原始作业活动获取数据，形成各产业漏洞管理水平和履责状态的“画像”，实现对全产业漏洞管理结果的“可视”、“可管”。尊重不同产业行业特点的差异，引导产业持续自我改进和优化，为客户提供更高效、更专业的漏洞管理支持。

5.4 应急响应

华为乾崮汽车安全应急响应中心SRC（Security Response Center）的职责是快速响应华为乾崮网络安全风险及问题，负责接受、处理和公开披露华为乾崮产品和方案相关的安全漏洞，同时也是公司对漏洞信息进行披露的唯一出口。



SRC遵循ISO/IEC 30111漏洞处理流程，以及ISO/IEC 29147漏洞披露标准，处理华为乾崮智能汽车解决方案中的漏洞。我们将按漏洞响应流程对上报的潜在漏洞进行处理。

漏洞响应流程介绍：

- 1、接受上报：**主动监控和接受外部上报安全漏洞和问题，启动漏洞响应流程。
- 2、问题验证：**协调资源，验证是否是漏洞或安全问题，评估风险等级。
- 3、解决方案：**制定漏洞风险缓解和修复方案。
- 4、漏洞披露：**漏洞确认修补后，与安全研究员协调安全问题的披露。
- 5、问题反馈：**收集和总结来自内外部客户的意见，重要案例反馈给开发流程和团队，用于指导产品开发。

为避免提前披露对消费者及行业合作伙伴造成伤害，在整个漏洞处理的过程中，我们会严格控制漏洞信息的范围，要求漏洞上报者对漏洞进行保密，直到漏洞修复和披露。

公司鼓励漏洞研究人员、行业组织、政府机构和供应商主动将与华为乾崮产品相关的安全漏洞报告给华为乾崮智能汽车解决方案汽车SRC组织。漏洞及网络安全问题上报渠道为：
psirt@yinwang.com；

华为乾崮智能汽车解决方案SRC PGP公钥：对于敏感信息，请使用此公钥PSIRT.asc加密后发送给我们。PGP fingerprint: 77F1 DB3E 2E65 B13D 100B F4ED BF81 BE78 7FC1 3523。

注意：请您给SRC邮箱发送加密邮件时，同时提供您的PGP公钥或链接，便于保证交互过程的安全性。



A 缩略语表

/ Acronyms and Abbreviations

表A-1 缩略语清单

缩写	英文全称	中文全称
APT	Advanced Persistent Threat	高级持续性威胁
AR-HUD	Augmented Reality Head-up Display	增强现实平视显示器
AutoSAR	Automotive Open System Architecture	汽车开放系统架构控
CAN总线	Controller Area Net	制器局域网
ECU	Electronic Control Unit	电子控制单元
GDPR	General Data Protection Regulation	通用数据保护条例硬
HSM	Hardware Security Module	件安全模块
LIN总线	Local Interconnect Network	低速串行总线
MCU	Micro Control Unit	微控制器单元
MDC	Mobile DataCenter	移动数据中心
OBD	On Board Diagnostics	车载自动诊断系统原
OEM	Original Equipment Manufacturer	型设备制造商
OTA	Over-the-Air Update	空中下载更新
PKI	Public Key Infrastructure	公共密钥基础设施重
RPMB	replay protected memory block	放保护存储区
SOC	System-On-a-Chip	片上系统
T-BOX	Telematics BOX	远程通信终端
TEE	Telecom Equipment Engineering	可信执行环境
TLS	Transport Layer Security	传输层安全协议
UDS	Unified Diagnostic Services	统一诊断服务
V2X	Vehicle To Everything	车联万物
VCU	Vehicle Control Unit	整车控制器
VIU	Vehicle Integrated Unit	网关连接和局部计算
VOS	Vehicle Control OS	智能车控OS