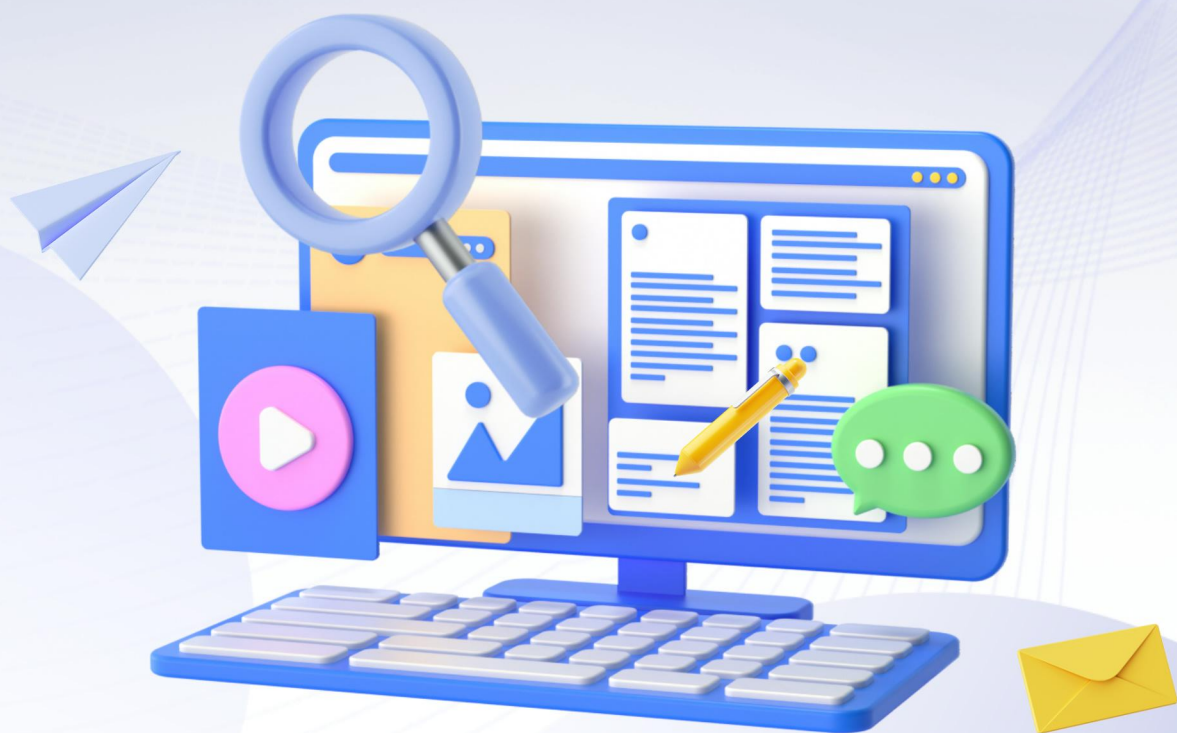




2023 年中国数据库行业 年度分析报告

墨天轮数据社区

2024 年 5 月

目 录

前 言	6
第一章：中国数据库发展现状	7
1、中国数据库流行度	7
2、中国数据库产品分析	11
3、数据库市场规模	16
4、金融核心系统替换进程	20
5、电信行业系统迁移改造	34
6、上市融资	36
7、开源生态	39
8、数据行业分析报告	44
9、数据库行业大事记	53
10、数据库国产化进程及标准化政策	74
11、数据库大会展现产业未来图景	78
12、数据库产业图谱	83
13、数据库社区年度发展概况	86
14、数据库高校学术力量蓬勃发展	95
15、墨天轮 2023 年度数据库获奖名单	100
第二章：数据库关键技术及发展趋势	110
1、云原生	110
2、HTAP	127
3、共享集群	131
4、超融合数据库	135
5、KV 数据库	142
6、多模态数据库	145
7、搜索型数据库	157

8、图数据库	166
9、文档数据库	182
10、时序数据库	188
11、向量数据库	195
12、流式数据库	201
13、空间数据库	203
14、数据仓库	207
15、数据湖	216
16、数据库兼容性	231
17、数据库内核创新	234
18、数据原生时代的数据库	241
19、从数据库到数据计算系统	244
20、自然语言交互	247
21、自动驾驶数据库	255
第三章：数据库新兴硬件及生态产品	259
1、处理器	259
2、存储	272
3、一体机	287
4、智能运维	294
5、云管平台	299
6、安全管控平台	301
7、SQL 审核	305
8、数据复制与集成	309
9、数据库中间件	312
总 结	317

表格目录

表 1	2023 年墨天轮排行榜年度 TOP 10 流行度趋势详情表	9
表 2	中国数据库流行度分类排名表	10
表 3	国产数据库在金融行业核心系统的分布情况	26
表 4	金融行业核心系统-国产数据库典型客户案例详情表	28
表 5	2023 年国产数据库融资情况统计表	36
表 6	2023 年国产数据库开源情况统计表	39
表 7	2023 年发布数据库重点评测	75
表 8	2023 年发布的数据库行业标准准则	77
表 9	2023 中国数据库厂商及社区年度大会一览	80

编委会成员

蔡冬者、杜剑峰、盖国强、洪春涛、洪建辉、黄向东、蒋锴、黎超、李文杰、梁敬彬、吕海波、栾小凡、罗小波、明玉琢、欧伟杰、潘巍、譙俊华、邵志杰、王振宇、魏可伟、魏兴华、吴疆、吴敏、萧少聪、衣国垒、俞炜婷、曾勇、詹年科、张晨、张俊成、张彤、章晨曦、章芋文、赵恒、赵新、周文雅、周颖冠、朱洁。

参编单位

爱可生、百度云、宝兰德、北京大学、贝格迈思、创邻科技、达梦数据库、Fabarta、飞轮科技、广东外语外贸大学、杭州电子科技大学、杭州图尔兹、华为技术有限公司、Intel、极限科技（INFINI Labs）、玖章算术、矩阵起源、巨杉数据库、科蓝软件、浪潮 KaiwuDB、蚂蚁 TuGraph、美创科技、PikiwiDB、清华大学、ScaleFlux、蜀天梦图、深算院 YashanDB、四维纵横、镜舟科技、拓数派、天谋科技、途普智能、沃趣科技、虚谷伟业、云尧科技、亚信安慧 AntDB、易景科技、云和恩墨、Zilliz。

（以上人物姓名和公司名称按首字母排序，排名不分先后）

版权声明

本报告著作权归墨天轮、各合作单位和个人共同享有，未经书面许可，任何机构或个人不得以任何形式翻版、复制、发表或引用。若征得墨天轮、各合作单位和个人同意进行引用、转载的，需在允许的范围内使用，并注明出处为“墨天轮”，且不得对本报告进行任何有悖原意的引用、删节或修改。本报告所涉及的观点或信息仅供参考，不构成任何投资建议。本报告的部分信息来源于公开资料，墨天轮、各合作单位和个人对该等信息的准确性、完整性或可靠性不做任何保证。本文所载的资料、意见及推测仅反映墨天轮于发布本报告当日的判断，过往报告中的描述不应作为日后的表现依据。在不同时期，墨天轮、各合作单位和个人可发出与本文所载资料、意见及推测不一致的报告和文章。墨天轮、各合作单位和个人不保证本报告所含信息保持在最新状态。同时，墨天轮、各合作单位和个人对本报告所含信息可在不发出通知的情形下做出修改，读者应当自行关注相应的更新或修改。

前言

在数字化浪潮的推动下，2023 年中国数据库行业站在了新的历史起点上。数据库作为信息化建设的核心，对于支撑数字经济的发展具有举足轻重的作用。随着云计算、大数据、人工智能等技术的飞速发展，数据库行业的技术革新和应用场景不断拓宽，市场需求日益增长，为国产数据库厂商带来了前所未有的发展机遇。

本年度报告旨在全面梳理 2023 年中国数据库行业的年度发展情况，从数据库流行度、市场动态、生态建设、资本投资，以及国产数据库在金融行业核心系统的替代进程等多个维度进行深入分析。同时，报告展现了数据库的技术演进，包括但不限于云原生数据库的兴起、分布式数据库的广泛应用、智能化数据库技术的发展趋势等。

2023 年中国数据库行业在技术创新、市场认可和国际竞争中取得了积极进展。国内厂商如 OceanBase、PingCAP、腾讯云和华为云等不仅在国内市场占据重要地位，还成功打入国际市场，赢得了广泛认可；在金融级分布式数据库领域，GoldenDB、GaussDB、OceanBase 和 TDSQL 等产品凭借出色的性能和安全性，成为市场的领导者。预计中国数据库市场将继续保持强劲的增长势头，本土厂商将进一步巩固和扩大其在全球数据库产业中的地位。特别是在金融、云服务和大数据等关键领域，中国数据库厂商有望发挥更大的影响力，推动整个行业向更高效、更安全、更智能的方向发展。

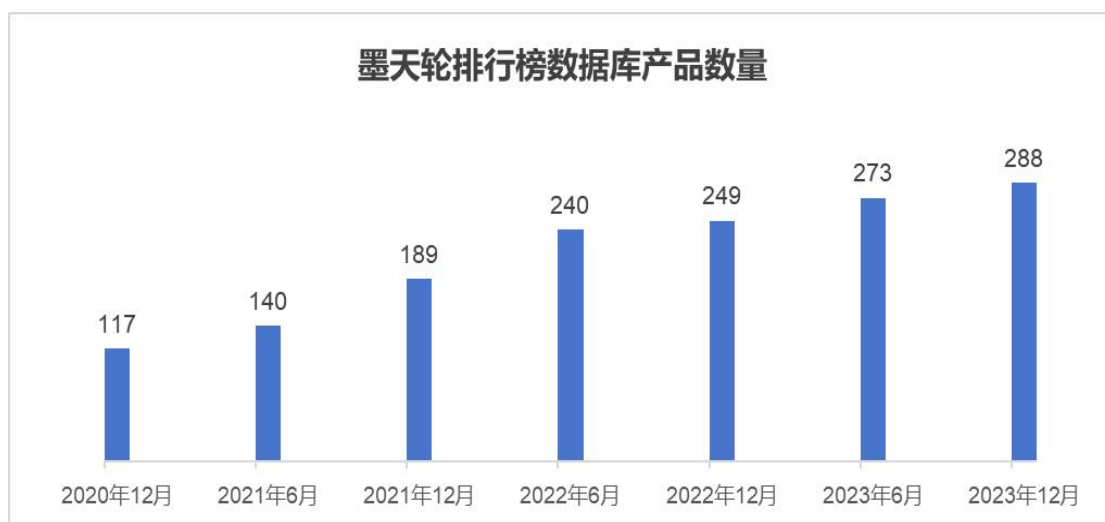
通过本报告，我们期望为业界提供一份客观、全面的行业发展参考，为数据库技术爱好者、企业决策者和专业人士提供宝贵的信息资源。我们相信，在政策支持、市场需求和技术进步的共同推动下，中国数据库行业将继续保持强劲的发展势头，在全球数据库领域占据更加重要的位置。未来，我们期待国产数据库能够继续深化技术创新，推动行业向更高层次的发展，为我国乃至全球的数字经济贡献更大的力量。

第一章：中国数据库发展现状

1、中国数据库流行度

一、中国数据库产品数量变化

墨天轮中国数据库流行度排行榜（简称：墨天轮排行榜）是一个专注于评估和展示中国数据库市场流行趋势的榜单。自 2019 年 6 月启动以来，该排行榜综合运用搜索引擎数据、趋势指数、资质认证、核心案例数、专利数、论文数、招聘数、书籍数量及平台优质内容数量等多个维度，每月更新一次，旨在客观反映中国数据库产品在互联网上的流行度和市场接受度。这一排行榜不仅为数据库用户提供了参考，也为国产数据库的发展提供了一个展示平台。



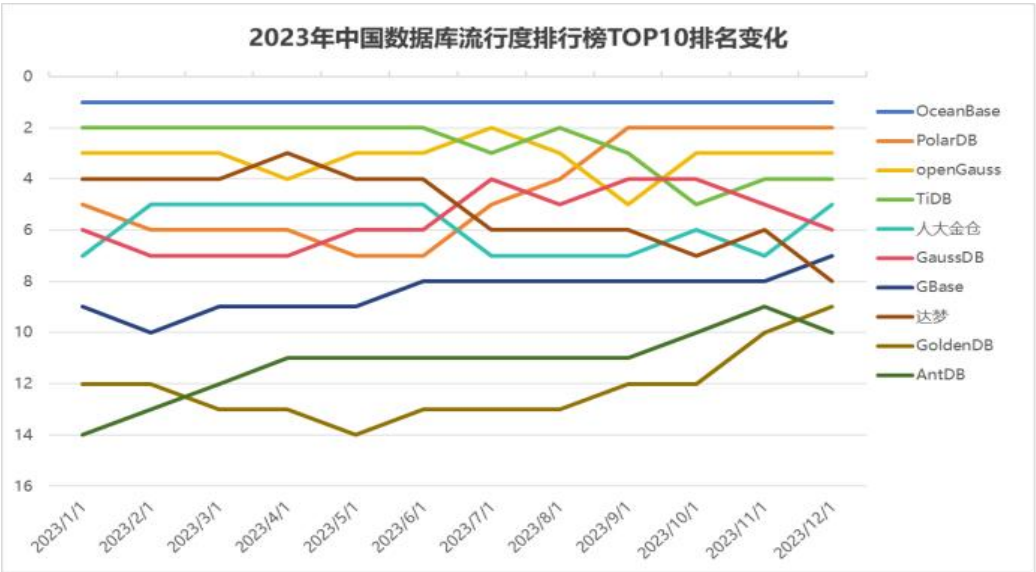
墨天轮排行榜数据库产品数量图

2023 年底，墨天轮排行榜收录了 **288** 个数据库产品，全年新增收录 39 个，相较于 2022 年的 55 款新增，数量减少了约 29%。新增数据库产品的数量减少，这反映出中国数据库市场在生产方向上正经历着由快速扩张向质量提升的转变。随着市场的逐步成熟，市场对于新入市者的技术和创新能力提出了更高的要求，同时也意味着现有竞争者需要在产品优化和服务升级上投入更多的努力，以维持和增强其市场地位。截至 2024 年 5 月，墨天轮排行榜共收录了 285 个数据库产品，收录数量的轻微下降是因为墨天轮排行榜在动态更新中，不断优化其收录标准，剔除了一些已停更或不再运营的数据库产品，以确保榜单的质量和时效性，为用户提供更为可靠和有价值的参考信息。

二、中国数据库流行度 TOP10

中国数据库行业在 2023 年继续保持着蓬勃的发展态势。墨天轮中国数据库流行度排行榜见证了这一年度的行业动态，各个中国数据库产品在互联网上的流行度和市场认可度得到了充分地体现。

2023 年中国数据库流行度排行榜年度 TOP10 依次为：OceanBase、PolarDB、openGauss、TiDB、人大金仓、GaussDB、GBase、达梦、GoldenDB 和 AntDB，这些数据库在技术创新和市场拓展方面均展现了不同程度的活跃态势和竞争力。



2023 年中国数据库排名前 10 名的竞争异常激烈。OceanBase 以其卓越的表现连续稳居排行榜首位，彰显了其强大的品牌影响力。第 2 名 - 第 4 名分别是阿里云 PolarDB、openGauss、TiDB，三者之间的得分差距较小，使得每次排名更新都充满悬念。GoldenDB 和 AntDB 是新加入 TOP10 的数据库，展现了不俗的竞争力。

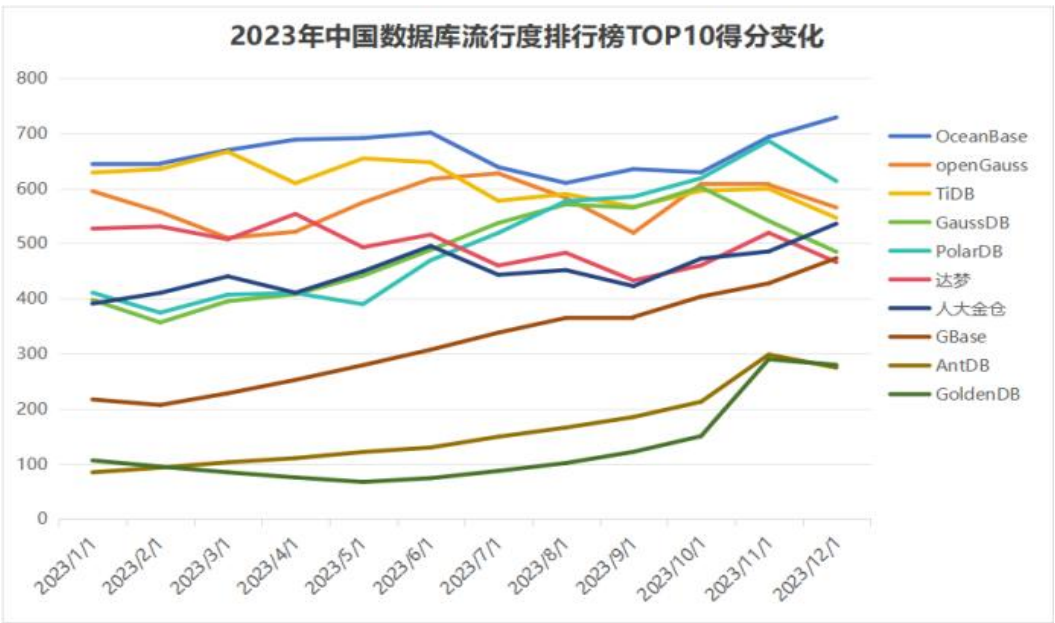


表 1 2023 年墨天轮排行榜年度 TOP 10 流行度趋势详情表

排名	数据库名称	主要趋势
1	OceanBase	稳居榜首，得分提升。OceanBase 全年稳居榜首，得分在年末达到最高，体现了其在市场中的领先地位和用户的高度认可，反映了其在网络上持续热度。
2	PolarDB	排名上升，得分大幅增长。PolarDB 的排名和得分均有显著提升，整体呈现稳步增长趋势，在第四季度稳住了第二的位置，显示了其在市场上的强劲动力。
3	openGauss	排名稳定，得分轻微下降。openGauss 的排名全年表现较为稳定，多次位列前三。得分有所波动但整体稳定，保持了其在市场上的高关注度。
4	TiDB	排名略有下降，得分减少。TiDB 作为往年的领先者，下半年排名波动较大，年底排名下降至第 4 名，但整体得分保持在高位，显示了其深厚的市场基础，在市场上的关注度依然很高。
5	GaussDB	排名稳定，得分增长。GaussDB 在年末的得分和排名均有显著提升，反映了其在市场上的快速成长和技术创新。
6	人大金仓	排名和得分均上升。人大金仓全年得分波动较大，但整体趋势向上，尤其在年末的排名中表现出色。
7	GBase	排名和得分显著提升。GBase 排名和得分均显著提升，成为年度成长最快的数据库之一。GBase 全年得分稳步上升，尽管在年中有轻微波动，但整体市场表现良好。
8	达梦	排名下降，得分减少。达梦数据库全年排名有所波动，整体排名有所下降，得分也有所减少，但在某些月份取得了较好的成绩，整体表现依然稳定。
9	GoldenDB	新晋 TOP10，表现稳定。GoldenDB 年末排名与得分提升较为明显，成功跻身年度 TOP10。其在技术性能、稳定性以及对金融级应用场景的深度适配上赢得了众多金融机构的信任和青睐，市场影响力提升。
10	AntDB	排名稳定，得分显著增长。作为新晋 TOP10，AntDB 整体表现依旧稳定，得分的持续增长彰显了其在市场上不懈努力和日益坚实的用户基础，预示着该数据库未来的发展潜力。

总结：

整体来看，这些数据库的排名和得分变化共同绘制了中国数据库市场一年来的动态图景，展现了国产数据库技术的进步和行业竞争的激烈态势。墨天轮中国数据库流行度排行榜为业界提供了一个观察和分析市场动态的窗口，也为数据库技术爱好者、企业决策者和专业人士提供了宝贵的参考信息。

三、中国数据库流行度分类排名

从数据库流行度来看，根据墨天轮中国数据库流行度排行榜，在关系型数据库领域，OceanBase、PolarDB、openGauss、TiDB 等国产数据库排名靠前。其中，OceanBase 连续 14 个月稳居榜首，阿里云 PolarDB 排名第二，openGauss 排名第三，TiDB 排名第四。时序数据库领域，涛思数据 TDengine 排名第一；向量数据库领域，爱可生 TensorDB 排名第一；键值数据库领域，腾讯云 TcaplusDB 排名第一；图数据库领域，北京大学 gStore 排名第一。

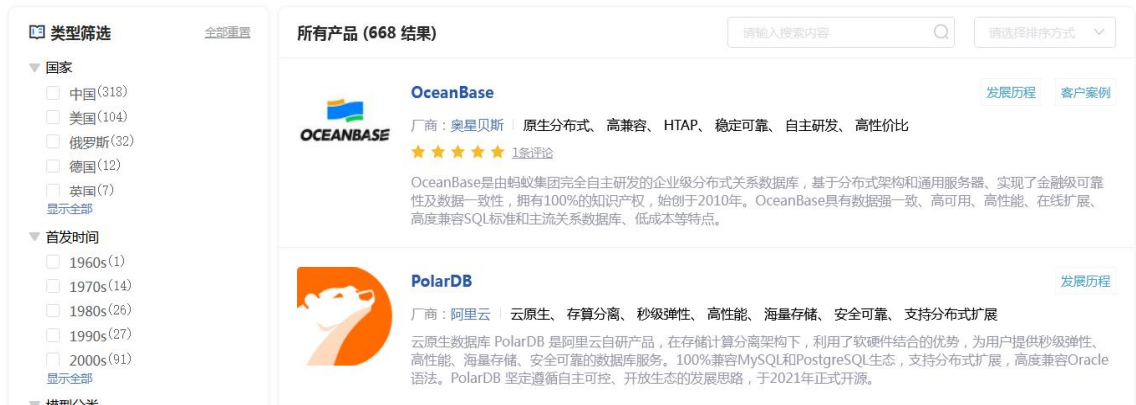
表 2 中国数据库流行度分类排名表

排名	关系型	时序	向量	键值	图
1	OceanBase	TDengine	TensorDB	TcaplusDB	gStore
2	PolarDB	KaiwuDB	腾讯云 VectorDB	Tendis	TuGraph
3	openGauss	DolphinDB	Milvus	Abase	NebulaGraph
4	TiDB	openGemini	cVector	KingDB	CSGGraph
5	人大金仓	IoTDB	DingoDB	Kvrocks	华为 GES
6	达梦数据库	CnosDB	Om-iBASE	ToplingDB	Galaxybase
7	GBASE	CTSDB	Vearch	FlashDB	TGDB
8	GaussDB	YMatrix	Hippo	TerarkDB	ArcGraph
9	TDSQL	NSDB	AwaDB	KeeWiDB	StellarDB
10	GoldenDB	GreptimeDB	ArcVector	Tair	阿里云 GDB

数据来源：墨天轮国产数据库流行度排行榜 2024 年 1 月数据

2、中国数据库产品分析

在 2023 年，墨天轮社区进一步丰富了其数据库百科功能，推出了全新的数据库产品分类统计功能。墨天轮的“数据库百科”产品功能提供了非常详尽的数据库产品信息，并且具有强大的类型筛选功能。这个功能允许用户根据不同的标准来筛选数据库产品，从而快速找到符合特定需求的数据库解决方案。目前墨天轮数据库百科收录了 **668** 款数据库产品。



墨天轮“数据库百科”产品功能

在类型筛选方面，墨天轮数据库百科提供了几百个维度的筛选选项：

- 国家：用户可以根据数据库产品的来源国家进行筛选，例如中国、美国、俄罗斯等。
- 成立时间：提供了数据库产品成立的年代，如 1960s、1970s 等，直到 2020s。
- 数据库类型：包括关系型数据库、键值数据库、时序数据库、图数据库、文档数据库等。
- 数据库起源：可筛选基于某些数据库研发的产品，如 PostgreSQL、MySQL、openGauss 等。
- 应用场景：例如 OLTP（在线事务处理）、OLAP（在线分析处理）、HTAP（混合事务/分析处理）等。
- 架构特点：如 Shared-Nothing、Shared-Everything、Embedded 等。
- 支持的功能：例如存储过程、认证体系、社区支持、Serverless 等。
- 测评认证：包括 ISO9001、ISO20000、ISO27001、Gartner 魔力象限、Gartner 市场份额、IDC 中国关系型数据库份额、沙利文数据库市场报告、IDC 数据库市场报告、Forrester Wave、CC EALA4+、CCRC EAL4+、分布式数据库金融标准验证、等保测评、ITSS、CMMI5、军 B+、GB18030、TPC-C、TPC-H、TPC-DS、可信数据库评测、国家安全可靠评测。
- 许可协议：如 Apache v2、Proprietary、GPL v2 等。
- 开发语言：涉及数据库开发所使用的编程语言，如 C++、Java、C 等。
- 索引类型：例如 B+Tree、Hash Table、Full-Text 等。
- 事务隔离级别：如 Read Committed、Serializable 等。

- 数据类型：支持的数据类型，如 JSON、DATE、CHAR 等。
- SQL 标准：支持的 SQL 标准版本，如 SQL-92、SQL:1999 等。
- 操作系统兼容性：数据库支持的操作系统，如 Linux、Windows、OS X 等。
- 硬件平台：数据库支持的硬件平台，如 Intel、鲲鹏、飞腾等。
- 编程语言接口：提供数据库访问接口的编程语言，如 Java、Python、C++等。
- 公司规模：数据库厂商的规模，如员工人数。
- 融资情况：数据库厂商的融资阶段。
- 总部城市：数据库厂商的公司总部位置。

这些筛选维度为用户提供了一个多角度、细致的数据库产品筛选机制，帮助用户从几百个不同的维度中找到最适合自己业务需求的数据库产品。通过这些筛选选项，用户可以更加精准地定位到特定的数据库产品，无论是从技术特性、应用场景、还是厂商背景等多个方面进行考量。

为了帮助用户更好地了解和应用国产数据库产品，墨天轮还推出了全新功能——“**国产数据库核心业务系统典型案例**”展示，旨在为广大用户提供一个展示国产数据库应用场景的窗口。平台已收录了大量基于国产数据库在企业核心业务系统的实战案例，这些案例涵盖了金融、电信、互联网、政务、制造业等多个领域。在这里，您可以按类型筛选来自不同行业、不同场景的国产数据库成功应用案例，还可筛选或搜索指定某个国产数据库的核心系统案例。目前，墨天轮“核心案例”功能，收录了 **488** 个国产数据库核心业务系统项目案例，其中精选案例 95 个。

案例筛选

全部重置

金融

一行二会 (3)

六大行 (18)

十二家股份制 (23)

保险 (39)

交易所 (9)

政策行 (4)

城商农信 (101)

证券 (46)

海外银行 (6)

基金 (24)

运营商

中国电信 (11)

中国移动 (90)

中国联通 (9)

中国广电 (2)

精选案例 (95 结果)

精选案例

请输入搜索内容

时间倒序

添加案例

审核案例

取消精选

编辑

删除

富邦华一银行

腾讯云TDSQL助力富邦华一银行新核心系统群全面上线

近日，富邦华一银行宣布新核心系统近期全面上线。该系统构建在国产分布式云平台腾讯云TCE之上，并采用国产分布式数据库腾讯云TDSQL承载核心系统数据，能够高效支撑富邦华一银行打造面向海量业务的分布式架构体系，为富邦华一银行新一代核心业务系统提供稳健支撑。富邦华一银行...

分布式数据库

案例详情

客户：富邦华一银行	数据库：TDSQL	投产时间：2023-12-31
系统：新核心系统	原数据库：无	场景：分布式数据库
原架构：无	新架构：无	容灾：无
采购方式：无	采购金额：无	性能提升：无

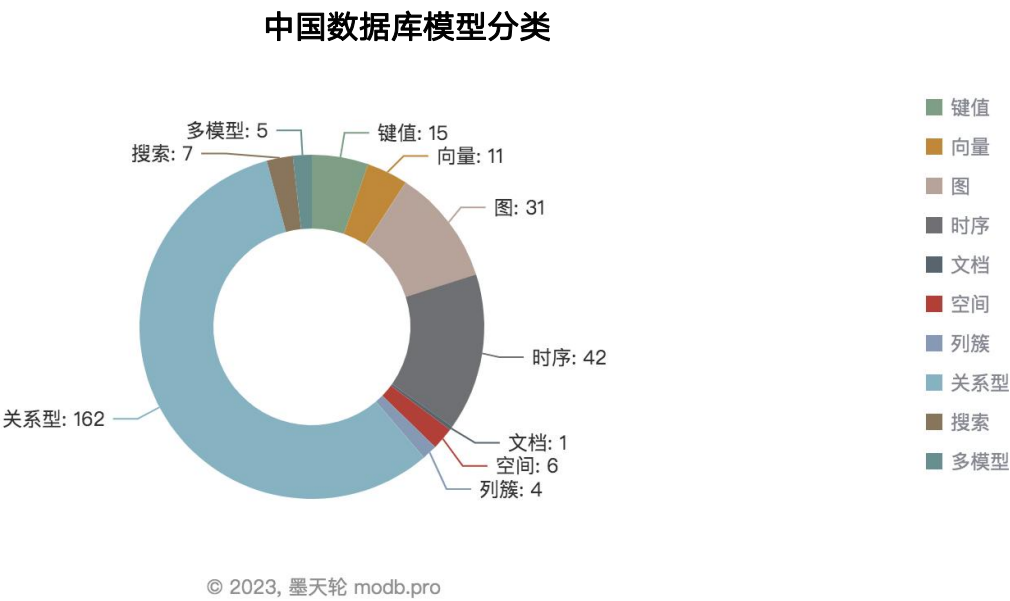
客户收益：腾讯云国产分布式数据库TDSQL支撑富邦华一银行打造了高可用的多活容灾核心系统数据库体系，提升富邦华一银行在推进大零售业务转型和服务小微与产业金融等发展目标时的系统数据处理能力和安全性。

墨天轮“精选核心案例”

通过具体案例的展示，不同行业的企业和机构可以了解到国产数据库在相似业务场景中的实践，为自己的技术选型和业务升级提供参考。

一、中国数据库模型

通过对中国数据库模型数量进行分析，关系型数据库以绝对优势占据主导地位，占比高达 57.0%，这反映了关系型数据库在数据管理和分析中的普遍适用性和重要性。时序数据库占比 14.8%位居第二，显示出市场对于处理时间序列数据的强烈需求，特别是在物联网、金融监控和实时分析等领域。图数据库占比 10.9%位列第三，凸显了图数据库在处理复杂关系和网络分析方面的重要性，其应用在社交网络分析、推荐系统等领域日益增多。

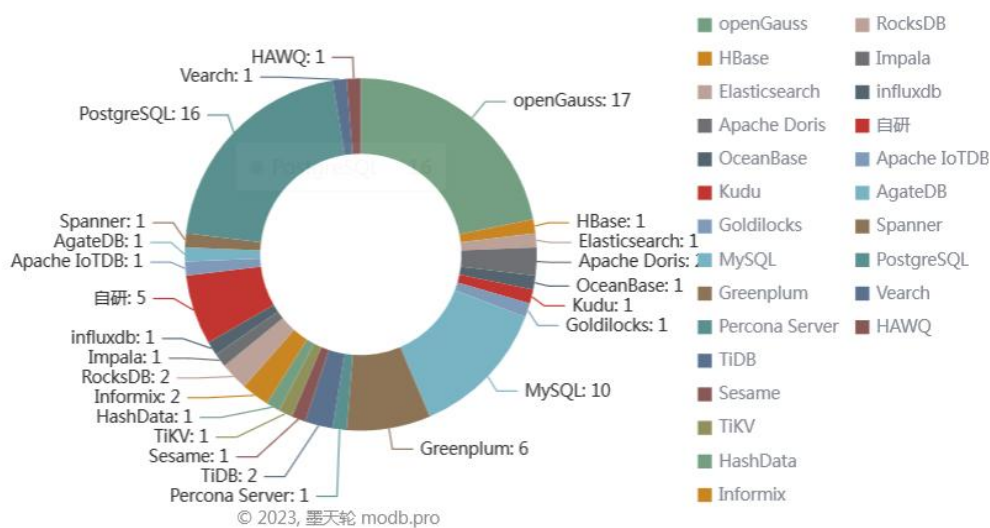


数据来源：墨天轮排行榜

二、中国数据库起源

据墨天轮对中国数据库起源的不完全统计，超 70 多款中国数据库产品是基于开源数据库开发的，起源数据库最多的前三名分别是 openGauss（占比 22.1%）、PostgreSQL（占比 20.8%）和 MySQL（占比 13.0%），其中，基于 openGauss 研发的中国数据库产品数量已与基于 PostgreSQL 研发的中国数据库产品数量相当。国产数据库产业在开源项目的基础上进行了广泛的创新和自主研发，形成了多样化的产品生态。国产数据库厂商在利用开源代码的同时，也在努力保持与开源社区的兼容性，并积极构建自己的生态系统，如 openGauss 社区正在吸引越来越多的企业参与。

中国数据库起源

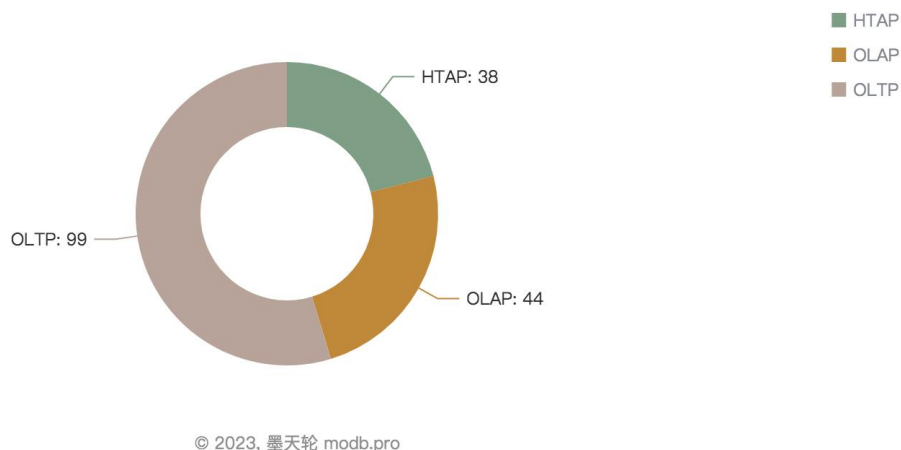


数据来源：墨天轮排行榜

三、中国数据库数据处理类型

在数据处理技术方面，HTAP、OLAP 和 OLTP 模型的数量占比分别为 21.0%、24.3%和 54.7%，反映了市场对实时数据处理和分析能力的高度重视。

中国数据库数据处理类型



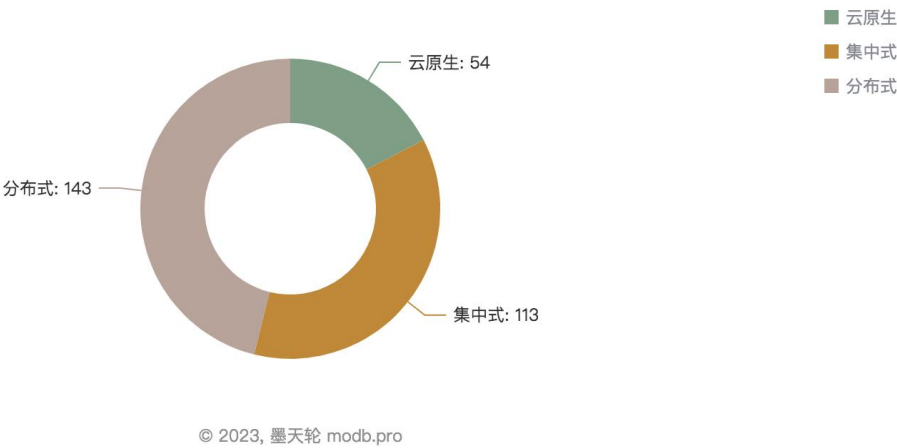
数据来源：墨天轮排行榜

四、中国数据库部署方式

随着云计算技术的不断进步，云原生和分布式数据库的部署方式逐渐成为主流。中国数据库部署方式上，云原生数据库占比 17.4%，而分布式数据库占比 46.1%，领先

于占比 36.5%的集中式数据库，这表明市场对于数据库的可扩展性和灵活性有着明确的偏好。

中国数据库部署方式

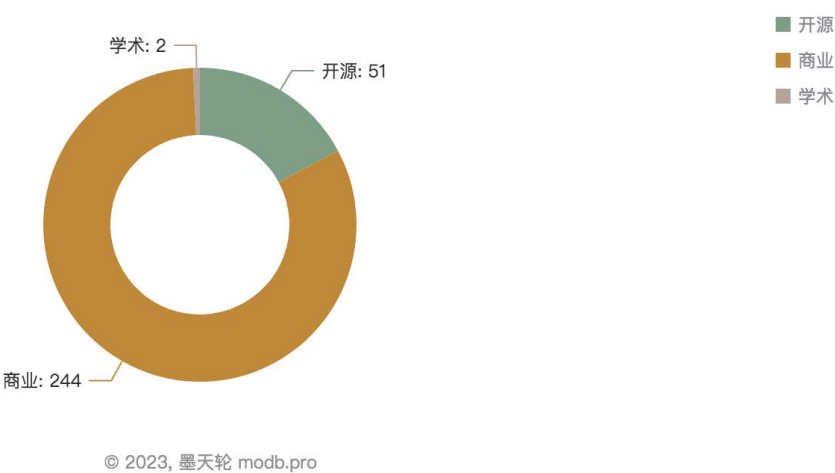


数据来源：墨天轮排行榜

五、中国数据库商业模式

在商业模式方面，中国商业数据库占比 82.2%主导着市场，这可能归因于其提供的专业性服务和全面解决方案。与此同时，开源数据库以 51 个模型的数量占比 17.2%，展现了开源技术在创新和社区支持方面的强大动力。

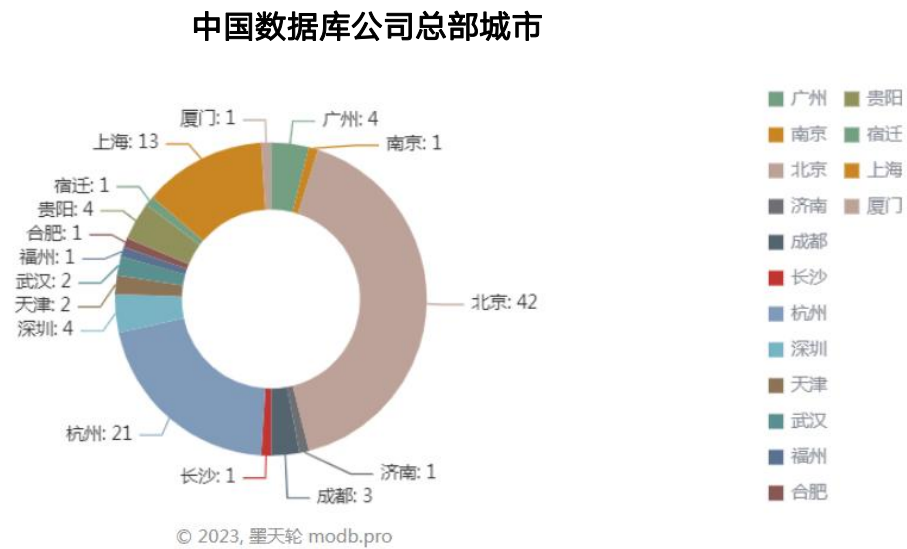
中国数据库商业模式



数据来源：墨天轮排行榜

六、中国数据库公司总部城市分布

北京作为中国数据库公司的集中地，总部数量占比 41.2%，位居首位。中国数据库公司总部在杭州和上海分别占比 20.6%和 12.7%，紧随其后。这些数据揭示了中国数据库产业在特定城市的集聚现象，显示了这些城市在政策、人才和资本方面的优势。



数据来源：墨天轮排行榜

总结：

2023 年，中国数据库行业展现出强劲的增长势头，其中云数据库和关系型数据库占据市场主导地位。墨天轮数据显示，关系型数据库保持市场核心地位，同时向量数据库的增长凸显了 AI 和机器学习技术对行业的影响。云原生和分布式数据库的部署方式逐渐取代传统集中式部署，展现出市场对灵活性和可扩展性的追求。北京、杭州和上海等城市成为数据库产业的集聚中心，显示了地域优势对产业发展的推动作用。总体上，技术创新和市场需求正驱动中国数据库市场向多样化和专业化发展，这将助力中国在全球数据库行业中占据更加重要的位置。

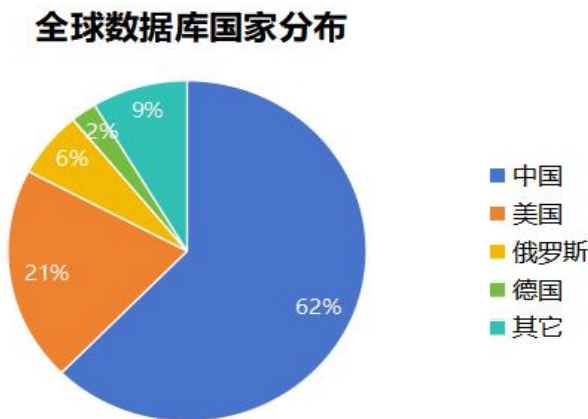
3、数据库市场规模

一、全球数据库市场

放眼全球市场，据 Gartner 预测，2023 年全球数据库规模将首次突破 1000 亿美元，其中，云数据库占比 55%。预计到 2027 年，云数据库的营收将占据全球数据库管理系统整体市场的 70%。这一数据不仅反映了数据库市场的快速增长，还强调了云

数据库相对于本地部署数据库的市场优势，预示着未来几年云数据库将继续保持强劲的增长势头。

据墨天轮网站统计，全球 668 款数据库产品中，中国数据库产品占全球数据库产品总数的 62%，**中国数据库产品数量全球第一**，其次是美国，美国数据库产品 104 个，占比 21%。中国和美国数据库产品数量的领先，表明了中美两国在数据库领域都有较强的研发能力和市场影响力。虽然中国数据库产品数量遥遥领先，但在全球数据库市场份额占比不高，也反映出中国数据库产品在市场竞争力、品牌影响力、产品成熟度等方面还有较大提升空间。



数据来源：墨天轮网站“数据库百科”统计

根据 DB-ENGINES 网站统计，目前市场上共有 400 多个数据库产品，其中关系型数据库的数量超过 160 个，约占总数的 40%。这一数据不仅展示了数据库产品的丰富多样性，同时也印证了关系型数据库在整体市场中的坚实地位和广泛适用性。

进一步聚焦中国市场，据墨天轮网站统计，在中国的 318 个数据库产品中，关系型数据库多达 182 个，占比高达 57%，在中国市场占据半壁江山。这一比例远高于全球平均水平，这也表明中国数据库产品在保持产品多样性的同时，关系型数据库依然占据着主导地位。

二、中国数据库市场规模

根据大数据技术标准推进委员会（CCSA TC601）测算，**2022 年全球数据库市场规模为 833 亿美元，中国数据库市场规模为 59.7 亿美元（约合 403.6 亿元人民币），占全球 7.2%。**预计到 2027 年，中国数据库市场总规模将达到 1286.8 亿元，市场年复合增长率（CAGR）为 26.1%。

2022—2027 年中国数据库市场规模及增速



来源：CCSA TC601，2023 年 6 月

按数据库部署方式划分市场规模，**2022 年中国公有云数据库市场规模为 219.15 亿元**，较 2021 年增速 51.6%，**本地部署数据库市场规模为 184.45 亿元**，较 2021 年增速 14.4%，公有云和本地部署模式市场规模分别占总市场 54.3%和 45.7%，2022 年公有云数据库市场规模首次过半，预计 2023 年公有云市场占比将进一步扩大达到 59.8%，规模达到 323.16 亿元，本地部署模式市场增速达到 17.8%，规模为 217.24 亿元。

2021-2023 中国公有云和本地部署数据库市场规模

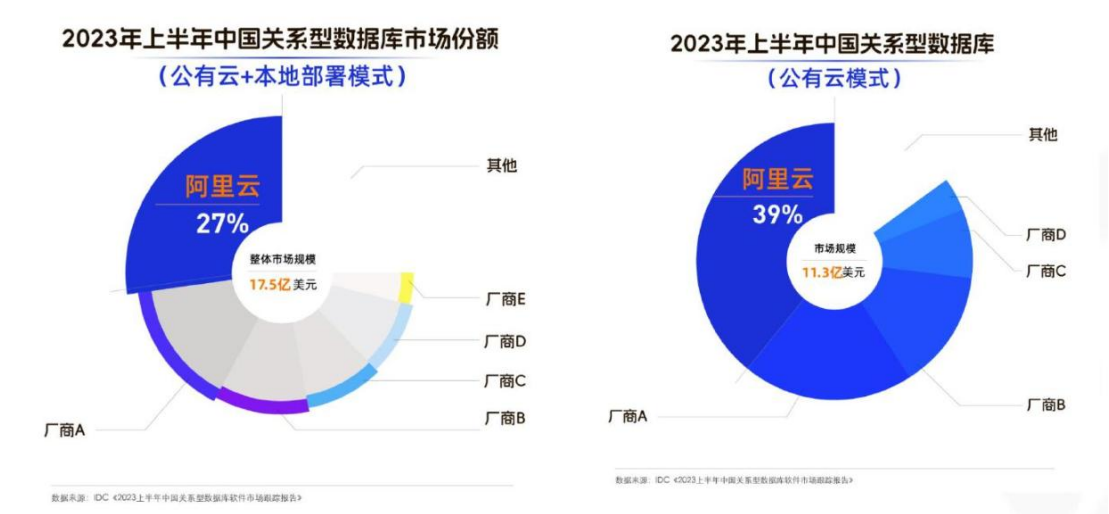


来源：CCSA TC601，2023 年 6 月

随着云服务商的数据库业务开始加大在私有云、行业客户等线下市场的发展力度，以及不断涌现的新兴关系型数据库厂商的入局，本地部署模式市场的竞争将愈发激烈。与本地部署市场相比，公有云关系型数据库的市场集中度更高。

根据 IDC 数据，2023 年在国内关系型数据库中，公有云数据库的市场中以阿里云、华为、腾讯为代表的国产数据库厂商份额已远超海外数据库厂商。

IDC 发布的《2023 年上半年中国关系型数据库软件市场跟踪报告》显示，2023 年上半年中国关系型数据库整体市场规模为 17.5 亿美元，同比增长 13%。其中，阿里云以 27% 的市场份额稳居第一，占比高于第二名和第三名的总和。



2023 年上半年中国关系型数据库市场份额图

2023 年上半年，关系型数据库公有云部署市场规模达 11.3 亿美元，同比增长 19%，市场规模是本地部署模式的 2 倍左右，选择公有云模式部署数据库已成为主流。在公有云部署模式中，阿里云占比 39%，以绝对优势蝉联第一。而在 2022 年下半年，中国关系型数据库公有云模式前五名企业包括阿里云、腾讯、AWS、华为、中国电信天翼云（前五名厂商份额共计 86.5%）。

2023 年上半年，华为云 GaussDB 数据库以 15.47% 的市场份额，位列本地部署模式国产厂商第一。2022 年下半年，中国关系型数据库市场份额本地部署前五名包括 Oracle、华为、微软、达梦、SAP。

随着企业数字化进程的推进和公有云数据库的优势日益凸显，公有云数据库业务作为“基础设施”有望核心受益，并向各个行业深度应用落地，其市场份额有望进一步扩大。

4、金融核心系统替换进程

一、背景概述

金融体系的稳定和安全直接关系到国家的经济稳定和社会发展，这不仅仅是个人和机构的利益问题，更是国家整体利益的重要组成部分。当前，金融服务已逐步扩展、细分，走进人们生活中的每一个场景，由此产生了海量数据和多样化的分析处理需求，而金融数据库在其中占据高度核心的地位，是保障金融业务顺利开展、稳定运行的关键基础。

而金融核心系统是处理和管理其最重要业务的关键系统，包括客户信息管理、支付结算、贷款管理、交易记录等，是金融机构运行的基石，由此对金融核心系统数据库的安全、稳定、可用、拓展等能力都提出了极高的要求。

在银行系统方面，数据库的应用具有以下几个显著特点：

1. 高可靠性：银行系统对数据的稳定性和安全性有着极高的要求，因此采用国产数据库可以满足其业务需求，确保数据不丢失。
2. 高性能：银行系统日常业务涉及大量的数据处理，因此需要具有高性能的数据库来应对高峰时段的并发请求。
3. 易于维护：银行系统数据库需要具备良好的可维护性，以降低运维成本和风险。
4. 兼容性：银行系统数据库需要与其他系统（如核心业务系统、中间件等）具有良好的兼容性，以确保业务的顺畅进行。

在保险系统方面，数据库的应用具有以下几个显著特点：

1. 数据量大：保险系统涉及众多业务数据，如保单、客户信息等，因此需要具备海量数据处理能力。
2. 实时性：保险系统中的部分业务（如理赔、销售等）对实时性要求较高，因此需要采用具备高并发处理能力的数据库。
3. 数据一致性：保险系统中的数据需要保持一致性，以确保业务正确性和客户利益。
4. 数据安全：保险系统涉及客户隐私，因此需要采用安全可靠的数据库技术来保护客户数据。

在证券系统方面，数据库的应用具有以下几个显著特点：

1. 高速处理：证券系统对数据处理速度要求较高，需要采用高性能的数据库来应对市场变化的快速反应。

2. 高并发：证券系统在交易时段需要处理大量并发请求，因此需要具备高并发处理能力的数据库。

3. 实时性：证券系统中的交易数据需要实时更新，以确保投资者获取准确的行情信息。

4. 数据安全：证券系统涉及资金安全和客户隐私，因此需要采用安全可靠的数据库技术来保护数据。

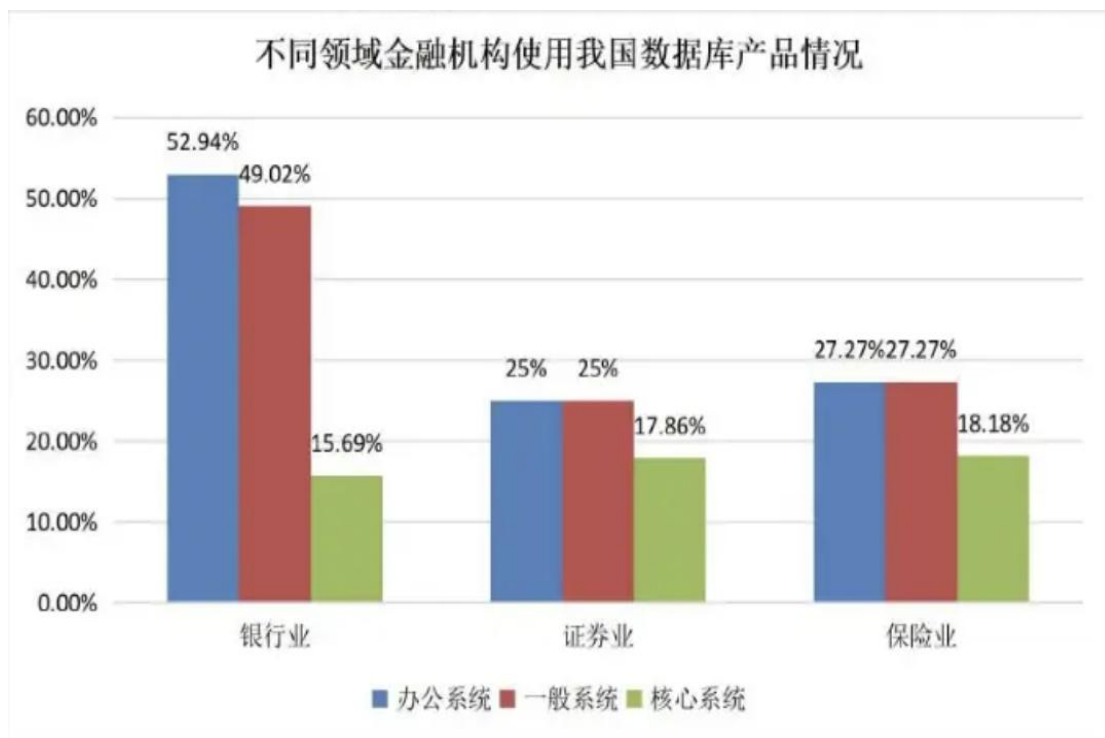
随着数字化转型不断深入，金融业务不断创新，对核心系统数据库的技术能力、功能多样化的需求逐步提高，原有数据库架构较难实现，各机构逐渐替换国产数据库系统。此外，在信创发展环境下，相关政策规划也推进了基础软件的自主可控，进一步加速了金融核心系统的数据库国产化落地。

当前，我国数据库厂商不断加强产品创新能力、提高产品硬实力，通过针对不同场景需求提供定制化解决方案，已在银行、证券、保险等多个金融行业的核心系统中得到了广泛应用。自 2019 年张家港农商行核心数据库首次实现国产化，接连中信银行信用卡 StarCard 新核心系统投产，这些大型金融机构的成功替换案例出现后，为行业打造了可复制、可参考的实践案例，此后金融核心数据库国产化进程逐步加快。

墨天轮针对部分国产数据库核心案例的不完全统计，**国产数据库已在中国银行、中国农业银行、招商银行、中国人保、国泰君安等金融机构核心系统共计成功投产 95 个**，范围涵盖央行、六大行、十二家股份制银行、城商农信等银行机构及大型保险、证券企业等。我国金融行业在核心系统数据库国产化方面已取得显著成果，由此直接反映中国数据库已经能够完全承载金融核心系统，且国产化进程正不断加速演进。

二、非核心系统国产数据库使用情况

过去，金融系统对国外数据库的依赖程度极高，核心业务系统几乎被国外厂商垄断，只在非核心业务的办公管理及一般系统中进行了国产替换。目前，国产数据库在金融行业非核心系统应用中取得了显著成效。根据金融信息化研究所 2022 年编撰的《金融业数据库供应链安全发展报告》中的数据显示，在银行，证券及保险等金融领域，办公系统和一般系统的使用占比均已达到较高比例。其中**在银行业，办公系统与一般系统，国产数据库使用比例均突破 50%**。金融行业整体超过 40% 的金融机构在办公和一般系统中实现了我国数据库产品的应用，成效明显。



数据来源：金融信息化研究所

图：不同领域金融机构使用我国数据库产品示意图

在这一背景下，达梦数据库、人大金仓和南大通用等国产数据库产品凭借其优秀的性能和本土化服务优势，在非核心系统的应用中展现出了强劲的竞争力。例如达梦数据库在山东农信办公自动化系统、中信证券综合办公平台的应用，全面保障了客户自动化办公数据安全，基于达梦数据库的中国铁建财务共享服务平台，有效提高了企业整体的财务水平。南大通用 GBase 8a MPP Cluster 凭借卓越的数据分析性能，在大数据平台、湖仓一体、实时数据仓库、数据集市、数据中台等场景提供稳定支撑，比如中国农业银行和中国银行的数据仓库项目、中国银行多家分行大数据平台项目、中国人保数据平台项目、山东省农信审计系统项目、江西银行和江苏银行的审计平台项目、证监会信息分析系统平台项目等。

除了在非核心系统应用中取得的成效，近期这些国产数据库也陆续在金融核心系统中得到应用，展现出国产数据库在关键技术上的突破和成熟。例如，达梦数据库在湖北银行和梅州客商银行的核心系统中实现了成功部署；金仓数据库支撑湘财证券完成 TA 系统国产化升级…… 这些成功案例展示了国产数据库在金融行业核心系统替代进程中的潜力和前景，为国产数据库在更广泛领域的应用奠定了坚实的基础。随着国产数据库技术的持续进步和金融行业对国产化替代需求的不断增长，预计国产数据库将在更多关键领域发挥更大的作用。

三、国产数据库早期金融自用场景

- OceanBase：支付宝、网商银行

支付宝自 2014 年双十一购物节起与 OceanBase 展开合作，面对 Oracle 集中式数据库在处理巨大流量和高并发时的不足，OceanBase 成功接管了支付宝 10% 的业务流量，保障了交易系统的稳定运行。到了 2016 年，支付宝已将所有核心数据，包括账务库，迁移至 OceanBase，该分布式数据库能够处理数十亿条 SQL，管理数百 PB 的数据量，并在超过百万的服务器核数上运行，全面支持支付宝的所有核心业务链路和五大业务板块，为支付宝提供了面对大规模流量和复杂业务场景时的数据库需求的可靠解决方案。



继支付宝之后，2015 年成立的全球首家云上银行——网商银行，于 2017 年全面采用 OceanBase 数据库，成为银行业中率先实现 100% 去 IOE 和自主可控的机构。OceanBase 帮助网商银行有效解决了数据可靠性与一致性问题，达到了 RPO 为零和 RTO 仅数分钟的高容灾标准，为金融数据安全建立了强有力的保障。这一合作不仅加强了网商银行在金融服务领域的竞争力，也为中国金融业在云计算和数据库技术方面的进步积累了宝贵的经验。



- 腾讯云 TDSQL：微众银行

2014 年微众银行成立，并与腾讯云开展合作，旨在将云计算技术应用于金融业务场景。2015 年 8 月，微众银行核心系统正式上线，TDSQL 也承载了微众银行数百个核心系统和全行所有 OLTP 业务。TDSQL 就此成为第一个使用在银行核心系统的国产数据库。这一过程中，腾讯云 TDSQL 从未出现过大规模系统故障或者数据安全问题，支撑了微众银行单日交易峰值近 6 亿笔，最高 TPS 达到 10 万+，为微粒贷、微业贷等业务的数百个核心系统提供坚实支撑，并极大降低户均 IT 成本。



四、银行核心业务系统替换元年

- 腾讯云 TDSQL：张家港农商银行

2019 年 9 月 16 日，张家港农商银行成功上线基于腾讯云 TDSQL 的新一代核心业务系统。这是银行首次采用国产分布式数据库，打破了对国外数据库的长期依赖，符合国家对金融核心领域技术自主可控的要求。新核心系统采用 x86 服务器，降低硬件成本至传统架构的 1/5 以下。性能表现卓越，高频账户类交易耗时 300 毫秒以内，查询类交易耗时 100 毫秒以内，批量代发代扣业务 20 秒内完成 1 万笔。TDSQL 支持在线横向扩展，解决性能瓶颈。



张家港农村商业银行

- GoldenDB：中信银行

中信银行于 2019 年 10 月 27 日启用了新的核心系统 StarCard，采用了由中信银行和中兴通讯联合研发的分布式数据库 GoldenDB。这是国内首个在大型股份制银行信用卡核心系统中成功应用的国产分布式数据库。目前 GoldenDB 支持多个业务系统，包括卡中心客户服务、营销支持、产品服务、信贷风险、运营支持等，使用了 X86 服务器和分布式集群解决方案（HBASE+ES+HIVE），实现了秒级时延的数据实时查询。



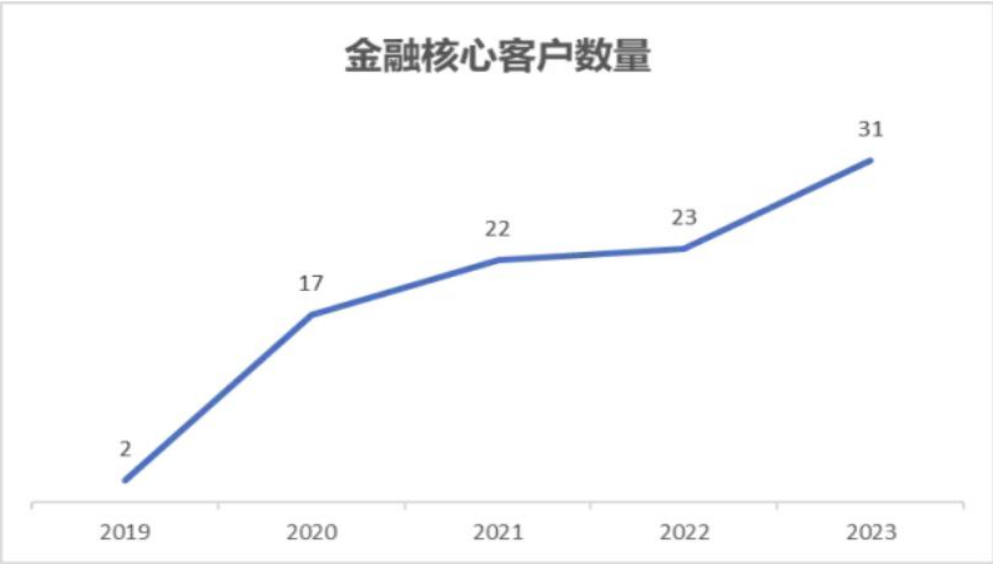
越来越多银行核心系统替换成功落地的案例，极大提升了业界对国产数据库的认可，也加速了替换的进程。同时，数据库系统需不断升级以适应不断涌现的业务需求，确保系统灵活性和可扩展性，数据库厂商也要提供更具成本效益的解决方案，以满足金融机构的财务需求。客户与厂商需共同努力，紧密合作，推动数据库技术的发展，以迎接行业变革带来的各种挑战。

五、金融核心系统国产化替换案例分析

（一） 国产数据库在金融行业核心系统应用范围扩大

在我国金融行业核心系统中，国产数据库的应用范围广泛，其中银行、证券和保险等行业均有涉及。墨天轮针对部分中国数据库产品的不完全统计，从 2019 年至 2023 年，金融行业核心系统采用国产数据库的公开典型客户案例共计 95 个。这些案例在数

量上呈现出逐年递增的趋势，表明了国产数据库在我国金融行业核心系统中的应用范围不断扩大，国产数据库助力金融行业系统国产化进程快速发展。

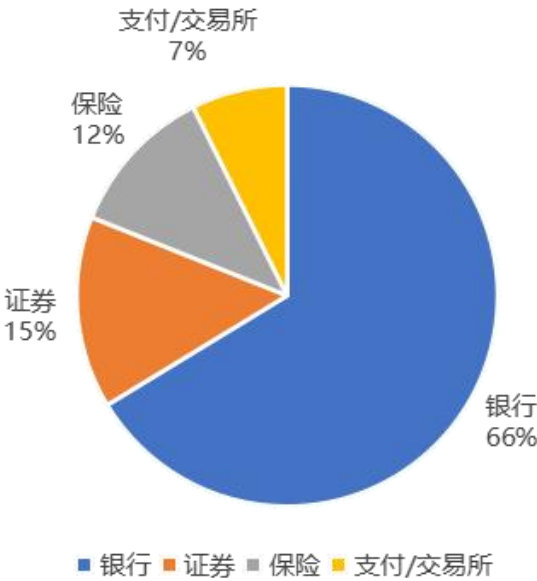


来源：墨天轮-数据库核心案例

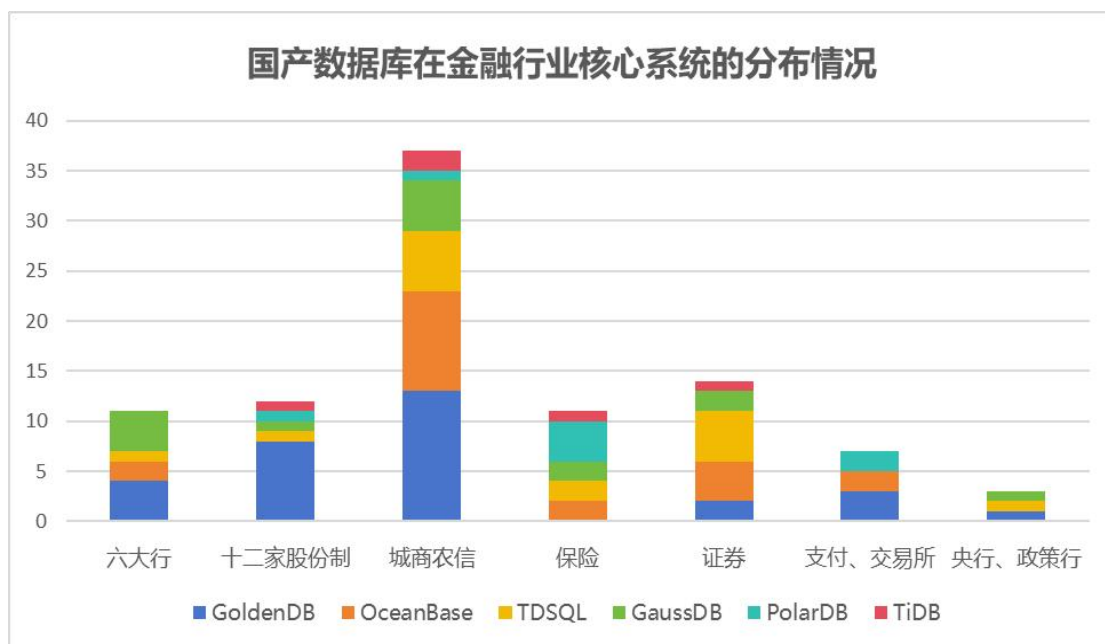
（二）国产数据库在金融行业核心系统的客户分布

根据已公开的金融行业核心系统-国产数据库典型客户案例数据统计，**银行客户占比达到 66%**，成为数据库应用的最大需求方。紧接着的是证券客户，占比为 15%，保险客户则位居第三，占比相对较低。

金融细分行业客户分布



来源：墨天轮-数据库核心案例



来源：墨天轮-数据库核心案例

表 3 国产数据库在金融行业核心系统的分布情况

数据库	六大行	十二家股份制	城商农信	保险	证券	支付、交易所	央行、政策行	合计
GoldenDB	4	8	13	/	2	3	1	31
OceanBase	2	/	10	2	4	2	/	20
TDSQL	1	1	6	2	5	/	1	16
GaussDB	4	1	5	2	2	/	1	15
PolarDB	/	1	1	4	/	2	/	8
TiDB	/	1	2	1	1	/	/	5
合计	11	12	37	11	14	7	3	95

来源：墨天轮-数据库核心案例

在银行案例中，金篆信科 GoldenDB 成为最受欢迎的数据库解决方案。据统计，GoldenDB 案例数量高达 26 个，覆盖了各类银行类型，包括六大行、十二家股份制银行以及城商农信银行等。具体来说，六大行中有 4 家采用了 GoldenDB，十二家股份制银行中有 8 家选择了该产品，城商农信银行中有 13 家予以应用。这些案例表明，GoldenDB 在银行领域的市场份额和认可度较高。

在保险行业中，阿里云 PolarDB 成为保险案例中应用最广泛的数据库，共有 4 家保险公司采用了该产品。PolarDB 在保险行业的应用逐渐成熟，为保险公司提供了稳定、高效的数据处理能力。

而在证券行业中，腾讯云 TDSQL 和 OceanBase 分别占据主导地位，分别有 5 家和 4 家证券公司使用。这些数据反映出，在保险和证券领域，阿里和腾讯云的数据库产品具有较高的市场份额。

综上所述，在我国金融行业核心系统中，银行客户对数据库的需求最大，证券和保险客户紧随其后。在众多数据库产品中，金篆信科 GoldenDB 在银行领域表现出色，阿里云 PolarDB 在保险行业占据一席之地，而腾讯云 TDSQL 和 OceanBase 则在证券市场展现出强大的竞争力。随着金融行业对数据库技术的不断需求和关注，这些数据库产品在未来有望取得更为广泛的应用。

在我国金融行业核心系统中，国产数据库的应用范围不断扩大，市场份额逐年上升。这充分证明了国产数据库在金融行业国产化替代进程中具有强大的生命力和竞争力。随着我国金融行业的持续发展，国产数据库在金融领域的应用将越发广泛，为推动我国金融信息化建设作出更大贡献。在此过程中，国产数据库厂商需不断提升技术研发能力和服务水平，满足金融行业日益严格的需求，助力我国金融行业走向更加繁荣。

在银行业中，国产数据库的应用已经取得了显著的成果。各大银行纷纷加大自主研发和创新的投入，致力于打造安全、高效、稳定的核心系统。以某大型银行为例，该银行在其核心业务系统中全面采用了国产数据库，不仅提高了系统性能，还大幅降低了运维成本。此外，国内多家中小银行也纷纷跟进，将国产数据库应用于核心业务系统，实现了业务流程的优化和信息技术的自主可控。

随着我国资本市场的快速发展，证券公司对于交易系统、风控系统等关键业务系统的需求日益增长。在此背景下，国产数据库凭借较高的性能和安全性，逐渐成为证券公司的首选。例如，某知名证券公司将其交易系统迁移至国产数据库，有效提高了系统稳定性和交易速度，为投资者提供了更好的交易体验。

此外，在保险行业，国产数据库同样表现出强大的竞争力。随着保险科技的崛起，保险公司对于核心业务系统、客户关系管理系统等关键应用的需求不断增加。国产数据库在保险行业的应用，不仅有助于提高业务处理速度，还能有效保障数据安全。某大型保险公司便将其核心业务系统迁移至国产数据库，实现了业务流程的简化、运营效率的提升以及客户服务的优化。

表 4 金融行业核心系统-国产数据库典型客户案例详情表

序号	分类	客户	数据库	上线时间	核心系统
1	六大行	建设银行	GaussDB	2020-11-15	建设银行信用卡核心系统
2	六大行	中国邮政储蓄银行	GaussDB	2022-04-23	邮储银行新一代个人业务分布式核心系统
3	六大行	中国工商银行	GaussDB	2022-07-25	工商银行信贷核心系统
4	六大行	中国农业银行	GaussDB	2023-06-01	超级网银核心系统、国际支付结算系统、监管报送系统
5	六大行	交通银行	GoldenDB	2020-08-13	交通银行信用卡中心
6	六大行	中国工商银行	GoldenDB	2021-01-01	IPVS 业务系统
7	六大行	中国建设银行	GoldenDB	2021-03-01	建设银行对私账务核心系统
8	六大行	中国银行	GoldenDB	2023-11-01	存款结算定
9	六大行	中国工商银行	OceanBase	2020-09-25	理财系统
10	六大行	交通银行	OceanBase	2022-01-01	贷记卡系统
11	六大行	中国农业银行	TDSQL	2022-03-01	中国农业银行客户信息和信用卡核心系统
12	十二家股份制	招商银行	GaussDB	2021-01-01	招商银行核心数仓平台
13	十二家股份制	中信银行	GoldenDB	2019-10-26	中信银行 StarCard 新核心系统
14	十二家股份制	浦发银行	GoldenDB	2020-08-01	浦发银行信用卡中心
15	十二家股份制	渤海银行	GoldenDB	2020-12-20	BCIF 业务系统
16	十二家股份制	民生银行	GoldenDB	2022-12-01	民生银行信用卡核心业务
17	十二家股份制	光大银行	GoldenDB	2023-03-01	光大银行信用卡核心业务系统
18	十二家股份制	广发银行	GoldenDB	2023-05-03	广发银行分布式银行核心系统
19	十二家股份制	恒丰银行	GoldenDB	2023-12-01	票据系统
20	十二家股份制	浙商银行	GoldenDB	2023-12-01	新一代票据系统

21	十二家股份制	渤海银行	PolarDB	2022-07-01	核心手机银行系统
22	十二家股份制	平安银行	TDSQL	2020-12-24	信用卡
23	十二家股份制	广发银行	TiDB	2023-11-03	客户信息系统
24	城商农信	广东省农村信用社联合社	GaussDB	2021-09-11	农信金融云
25	城商农信	江南农村商业银行	GaussDB	2023-06-15	信贷核算核心系统
26	城商农信	华夏银行	GaussDB	2023-07-05	华夏银行核心业务系统(借记卡)
27	城商农信	江苏银行	GaussDB	2023-09-03	网贷系统、用户中台等
28	城商农信	苏州农村商业银行	GaussDB	2023-09-07	超级网银核心系统
29	城商农信	赣州银行	GoldenDB	2020-10-01	新一代信贷系统
30	城商农信	贵州银行	GoldenDB	2020-11-16	新一代信息系统
31	城商农信	广州银行	GoldenDB	2021-02-01	大额现金业务系统,视频银行,房抵贷
32	城商农信	深圳农村商业银行	GoldenDB	2021-07-01	深圳农商行互联网金融核心业务平台
33	城商农信	南京银行	GoldenDB	2021-12-01	南京银行理财核心业务系统
34	城商农信	徽商银行	GoldenDB	2021-12-01	ECIF、对公特色业务、对公特色业务、互联网核心
35	城商农信	云南农信	GoldenDB	2021-12-01	统一信息发布平台
36	城商农信	广东农信	GoldenDB	2021-12-01	新一代柜面系统, 上账务核心
37	城商农信	江苏农信	GoldenDB	2023-07-01	信贷系统
38	城商农信	山西银行	GoldenDB	2023-10-01	业务中台、智慧网点平台、互联网渠道平台
39	城商农信	山东省城商联盟	GoldenDB	2023-11-01	电子银行、网银互联、支付平台
40	城商农信	甘肃银行	GoldenDB	2023-11-01	信用卡核心
41	城商农信	齐鲁银行	GoldenDB	2023-12-01	零售信贷业务系统

42	城商农信	顺德农村商业银行	OceanBase	2020-07-06	顺德农商银行核心系统
43	城商农信	云南红塔银行	OceanBase	2021-05-29	新核心
44	城商农信	西安银行	OceanBase	2021-10-22	西安银行互联网金融业务平台
45	城商农信	富滇银行	OceanBase	2022-01-16	电子银行,智慧信贷,风控平台,运营平台
46	城商农信	常熟农商银行	OceanBase	2022-04-10	核心业务
47	城商农信	四川农信	OceanBase	2023-01-10	四川农信“蜀信云”平台
48	城商农信	金华银行	OceanBase	2023-06-12	金华银行新一代核心系统
49	城商农信	北京银行	OceanBase	2023-08-29	北京银行客户信息系统
50	城商农信	深圳农商银行	OceanBase	2023-09-01	银行卡前置系统,综合理财投资系统
51	城商农信	宁夏银行	OceanBase	2023-12-01	信贷,移动作业平台,外部数据管理
52	城商农信	中和农信	PolarDB	2023-04-01	信贷业务系统
53	城商农信	张家港农商	TDSQL	2019-08-16	张家港行新一代核心系统
54	城商农信	昆山农商银行	TDSQL	2021-08-08	昆山农商行新一代核心系统
55	城商农信	福建海峡银行	TDSQL	2022-04-12	福建海峡银行新一代核心业务系统
56	城商农信	秦皇岛银行	TDSQL	2022-07-10	TDSQL 秦皇岛银行分布式核心系统
57	城商农信	湖州银行	TDSQL	2023-05-04	新核心
58	城商农信	富邦华一银行	TDSQL	2023-12-31	新核心系统
59	城商农信	北京银行	TiDB	2020-12-01	网联支付清算平台,银联无卡快捷支付系统,手机银行
60	城商农信	杭州银行	TiDB	2023-01-02	杭州银行新一代分布式核心系统
61	央行	中国人民银行清算总中心	GaussDB	2021-07-23	大额支付平台和超级网银系统

62	央行	中国人民银行清算总中心	TDSQL	2021-08-19	支付系统 PQDB 国产化系统
63	政策行	国家开发银行	GoldenDB	2022-05-01	国家开发银行新一代核心系统
64	保险	中国人保	GaussDB	2022-07-15	健康险、寿险等核心系统
65	保险	永安保险	GaussDB	2022-10-28	永安保险核心业务系统
66	保险	中华财险	OceanBase	2021-02-09	新核心
67	保险	太平洋保险(集团)股份有限公司	OceanBase	2022-12-18	太保 P17 核心系统
68	保险	横琴人寿	PolarDB	2020-10-22	寿险互联平台
69	保险	友邦人寿	PolarDB	2021-08-04	保险业务系统
70	保险	中国人寿保险股份有限公司	PolarDB	2021-11-04	金融核心系统
71	保险	弘康人寿	PolarDB	2022-11-19	核心业务系统
72	保险	民生保险	TDSQL	2022-05-12	新一代寿险核心业务系统
73	保险	中国太平	TDSQL	2022-12-01	核心系统
74	保险	中国人寿财险	TiDB	2020-11-17	中国人寿财险核心业务系统
75	证券	国信证券	GaussDB	2022-12-01	集中运营平台,开放式基金系统,金太阳手机证券系统
76	证券	兴业证券	GaussDB	2023-03-21	兴业证券新一代法人清算系统
77	证券	国泰君安	GoldenDB	2022-08-22	新一代分布式低延时交易平台
78	证券	中信建投	GoldenDB	2023-09-01	极速快交
79	证券	招商证券	OceanBase	2020-12-09	历史收益系统
80	证券	广发证券	OceanBase	2021-12-01	新一代估值系统
81	证券	方正证券	OceanBase	2021-12-01	新一代认证系统

82	证券	安信证券	OceanBase	2023-01-10	资管清算系统
83	证券	东吴证券	TDSQL	2020-06-29	东吴证券核心交易系统
84	证券	富途证券	TDSQL	2020-08-06	富途核心系统
85	证券	中信建投	TDSQL	2021-01-01	全历史客户分析系统,交易系统
86	证券	中信信托	TDSQL	2021-01-01	营销管理平台,消费金融系统,APP
87	证券	国信证券	TDSQL	2022-01-01	国信证券新一代 OTC 交易柜台系统
88	证券	安信证券	TiDB	2022-12-01	客户资产中心系统,ATP 极速交易系统
89	支付	银联数据服务有限公司	GoldenDB	2020-06-29	银联数据新一代信用卡系统
90	支付	信达资产	GoldenDB	2023-05-01	信达资产核心业务系统
91	支付	网商银行	OceanBase	2020-12-21	核心系统
92	支付	汇付天下	PolarDB	2022-10-14	核心业务系统
93	支付	富友支付	PolarDB	2023-08-11	核心扫码交易系统
94	交易所	深圳证券交易所	GoldenDB	2022-01-01	集中交易系统
95	交易所	郑州商品交易所	OceanBase	2023-02-08	郑州商品交易所业务系统

来源：墨天轮“数据库核心案例”精选案例

（三）国产数据库进军海外市场

此外，我国国产数据库在海外市场，已经取得了一定的成绩。特别是在金融领域，特别是在金融领域，已逐渐赢得了海外客户的信任和认可。这些海外金融客户主要为阿里、腾讯等我国知名大型企业所有。

● OceanBase：印度尼西亚电子钱包 DANA

蚂蚁旗下的自研数据库 OceanBase 在印度尼西亚电子钱包 DANA 项目中发挥了重要作用，为用户提供稳定、高效的数据服务。同样，在菲律宾国钱包 APP GCash 项目中，OceanBase 也展现了强大的技术实力。此外，它还在非洲推出了“支付宝” PalmPayWallet 业务，为当地用户提供便捷的支付服务。



- **阿里云 PolarDB：东南亚金融科技公司 Akulaku**

阿里云 PolarDB 在东南亚金融科技公司 Akulaku 的案例中，为该公司提供了高性能、可扩展的数据库解决方案，助力其业务快速发展。迁移到 PolarDB 后，主从同步的效率有明显地提升。



- **腾讯云 TDSQL：印尼 BNC 银行**

腾讯云数据库 TDSQL 助力印尼 BNC 银行完成新核心分布式迁移，告别昂贵的传统商业数据库，实现数字化转型。新系统运行平稳顺畅，截至 2022 年底，已支持超 2000 万用户每月支撑月活用户 50%以上的增长，每天处理 200 万笔交易和 15 万笔贷款发放，完成超过 9 亿美元月交易额。



这些成功案例充分展示了国产数据库的技术实力和潜力，相信在未来，它们将在全球范围内发挥更加重要的作用。在这个过程中，阿里、腾讯等大型企业将继续发挥引领作用，推动国产数据库走向世界舞台。我们期待在未来，国产数据库在全球市场上取得更加辉煌的成就，为全球用户带来更多便捷、高效的数据服务。

总结：

当前在我国金融行业核心系统中，国产数据库的应用已经取得了显著的成果，且正处于持续高速发展阶段。历经真实场景的不断打磨，国产数据库已具备承载金融核心系统的能力。随着金融行业的发展和国产数据库技术的不断创新，相信未来国产数据库在金融领域的应用将进一步拓展，在金融核心系统中的落地投产数也将进一步攀升。同时，借着信创产业国产替代的东风，国产数据库更将为我国金融系统的稳定运行和创新发展提供有力支持。

5、电信行业系统迁移改造

今天，大中型企业在 IT 方面遇到的问题并不罕见，电信运营商在几年前就遇到过。BOSS 系统（业务运营支撑系统）是电信行业处理业务信息、进行集约化管理的重要基础平台，以客户服务、业务运营和管理为核心，以客户服务和计费为重点和关键性事务操作，为运营商企业提供综合的业务运营和管理服务，而这与各行业企业实现数智化转型的道路几乎相同。

随着自主可控战略的不断深入，电信行业核心系统的替换成为必然趋势。这种趋势对数据库行业产生了深远的影响，因为数据库在支持电信行业系统的性能和稳定性方面，发挥着至关重要的作用。

对于电信运营商来说，传统的支撑系统包括了 M 域（管理域）、B 域（业务域）和 O 域（网络域）。通信运营支撑系统（Business & Operation Support System），涵盖了以往的计费、结算、营业、账务和客户服务等系统的功能，对各种业务功能进行集中、统一的规划和整合。电信行业 BOSS 系统架构对于需要引入数智化转型赋能的企业来说，具有重要参考意义，如下图。



图 1：电信行业支撑系统示意图

数据库迁移改造是一项复杂的工程，不仅是将数据从一个地方移动到另一个地方，还需要考虑业务定义、架构变更、应用改造、数据安全等诸多方面问题。

CRM 全域系统涵盖客户、订单、产商品中心等多个关键领域，涉及市场营销、销售实现、客户服务等多个领域，有着规模大、业务多、流程复杂等特点。BOSS 核心计

费系统包含计费、账务、账管在内，承载着数千万用户的充值缴费、账务记录的工作，业务处理量大，是业务支撑系统转型的关键点之一，对数据库的可靠性、性能要求极高。

在迁移策略上，CRM 域采用“先边缘，后核心”策略逐步全域自主可控替换；BOSS 域采用关系型数据库存放资金、账单、资料等，内存数据库存放高并发热数据的方式。

在突破“国内首个”BOSS 域、某省 CRM 全域核心数据库迁移改造过程中，目前国内数据库不仅实现了与原有数据库硬件的一比一替换，还展现出了国内数据库在同等硬件情况下的极强可用性；在内存数据库的高效性、分布式数据库的可扩展性等方面也有突破性进展。分布式数据库具备的数据节点在线动态扩展能力，能够有效解决数据暴涨情况下的“容量焦虑”问题。

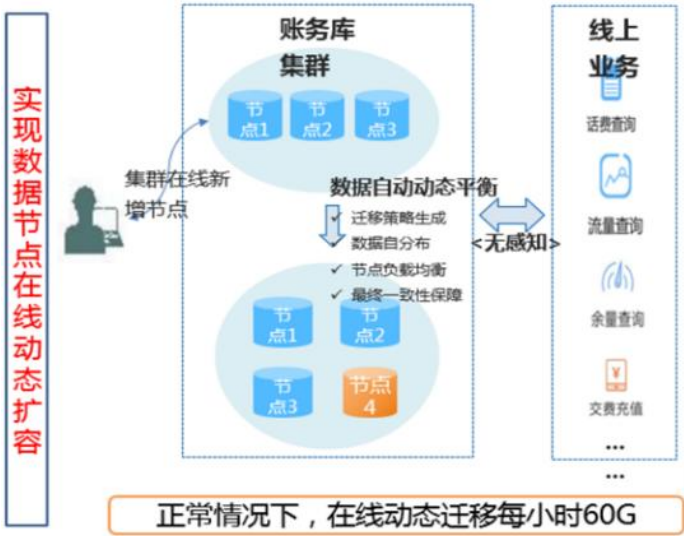


图 2：快速地在在线扩展能力

核心系统数据库的迁移改造，不仅要有合理的数据库选型；也要结合系统特点，制定针对性的改造方案；同时，综合考虑数据库的运维特性，制定合适的应对措施，降低或者消除风险，保障数据库迁移改造的顺利进行。

迁移流程如下：

- （1）初步评估：使用迁移评估工具（如：AntDB-MTK）对要迁移的数据库进行评估兼容性情况，并给出评估报告；
- （2）业务改造：对于极少量不兼容的场景，业务需要进行调整，并且回归测试；
- （3）评估容量：迁移人员需要评估数据库的建设和规划目标，包括数据容量等；
- （4）数据迁移：使用数据库同步工具（AntDB-MTK 等）进行数据和源数据迁移。
- （5）数据校验：对迁移前和迁移后的数据，进行一致性校验。

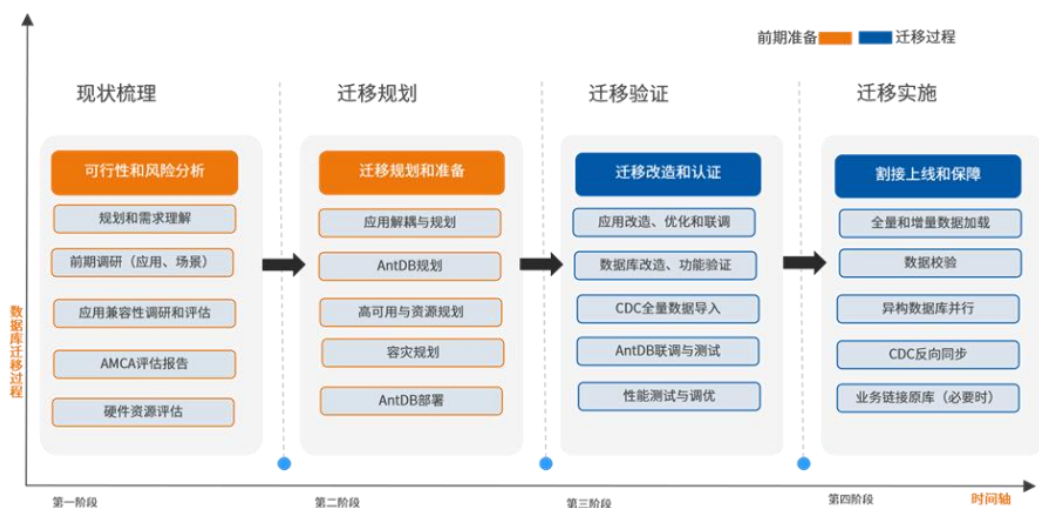


图 3：迁移改造整体计划流程图

6、上市融资

一、国产数据库领域投资降温

据不完全统计，2023 年仅有 10 家国产数据库厂商成功获得融资。相较于去年的 17 家，成功融资的国产数据库厂商数量下滑明显，明显反映出国产数据库投资热潮的冷却。2023 年融资主要集中在初创厂商，其中包括九章云极、四维纵横等企业，它们分别获得了超过亿元人民币的融资。

表 5 2023 年国产数据库融资情况统计表

融资时间	公司名称	轮次	金额	成立时间	投资方	数据库
2023/2/16	易鲸捷	股权融资	6515 万	2015/12/15	中国软件、中国电子、中国物流集团	QianBase
2023/4/6	泽拓科技	A 轮	未披露	2020/12/21	复星创富，老股东常春藤资本	Klustron
2023/5/24	航天软件	IPO 上市	12.68 亿人民币	2020/12/12	/	神通
2023/6/1	飞轮科技	Pre-A 轮	数亿元人民币	2021/12/1	IDG 资本、红杉中国等	SelectDB
2023/6/6	力控科技	B 轮	未披露	2011/4/12	中国石化资本	pSpace PLus
2023/7/13	四维纵横	A+轮	未披露	2020/8/24	用友网络、龙辕资本	YMatrix

2023/9/28	第四范式	IPO 上市	8.36 亿港元	2014/9/17	/	OpenMLDB
2023/10/8	九章云极	D 轮	3 亿人民币	2013/2/26	中电智慧基金、华民基金、中国太平、东方嘉富、卓源资本、中文在线	DingoDB
2023/10/19	四维纵横	B 轮	超亿人民币	2020/8/24	用友网络、顺义产业、龙辕资本	YMatrix
2023/11/23	沃趣科技	B+轮	数千万人民币	2012/4/28	翌马投资（恒生电子产业基金）、金蚂投资	QData

二、2023 年国产数据库融资事件

4 月 6 日，分布式 HTAP 数据库初创企业“泽拓科技”宣布完成 A 轮融资，投资方为复星创富，老股东常春藤资本持续加码。泽拓团队基于各核心主创在多个公有云大厂打造多款主力数据库产品的经验，目标是研发一款高性能，高可用的金融级分布式 HTAP 数据库产品 - Klustron，服务全球客户。

6 月 1 日，北京飞轮数据科技有限公司完成了 Pre-A 轮数亿元融资。飞轮科技是一家基于 Apache Doris 的商业化公司，2022 年 1 月成立，4 月完成天使轮和天使+轮融资，由 IDG 资本、红杉中国等顶级 VC 投资，融资金额超过 3 亿元。飞轮科技为 Apache Doris 用户提供技术支持商业服务，并推出了推出基于 Apache Doris 内核的 SelectDB 商业产品。



6 月 6 日，实时历史数据库 pSpace 厂商北京力控元通科技有限公司获得中国石化集团资本有限公司战略融资。本轮股权融资完成后，力控科技持续打造工业软件和工控信息安全的基础平台及适配产品，更好地服务于广大客户。



10月8日，九章云极 DataCanvas 完成总融资额 3 亿元 D1 轮融资。中国电子集团旗下中电智慧基金、华民投、中国太平旗下太平创新、浙江东方旗下东方嘉富等央企旗下投资机构，以及卓源资本等专注人工智能赛道的知名财务投资机构参与本轮融资。



10月19日，新一代超融合数据库厂商“四维纵横”完成超亿元人民币的 B 轮融资，本轮融资由用友、顺义产业基金领投，广州同创基金跟投。本轮融资资金将主要用于核心技术研发与商业生态建设，进一步提升 YMatrix 超融合数据库在全场景下的性能与生态优势。



10月20日消息，近日，北京柏睿数据科技股份有限公司上市辅导备案获北京证监局受理登记，辅导机构为安信证券。柏睿数据启动 IPO 后，国产数据库阵营或将出现又一家上市公司。



11月23日，沃趣科技完成了数千万 B+轮战略融资。本轮融资由翌马投资（恒生电子产业基金）领投，金蚂投资跟投，融资金额将用于研发投入及市场拓展。



12月20日，武汉达梦数据库公司的科创板上市，已获得证监会批准。资料显示，达梦数据本次 IPO 融资规模为 23.51 亿元，将用于集群数据库管理系统升级项目、高性能分布式关系数据库管理系统升级项目、新一代云数据库产品建设项目、达梦中国数据库产业基地以及达梦研究院建设项目。



总结：

尽管今年国产数据库融资整体呈现下滑趋势，但初创厂商仍备受投资方青睐。2023 年国产数据库融资情况反映出投资市场的冷静思考。在经历了过去几年的快速发展后，投资者对国产数据库产业的预期变得更加理性和谨慎。尽管如此，仍有部分优秀初创企业成功获得了投资，显示出我国数据库产业的活力和潜力。在政策支持和市场需求的双重推动下，我们相信国产数据库厂商将继续加大创新力度，提升产品竞争力，为我国数据库产业的繁荣做出更大贡献。

7、开源生态

一、2023 年国产数据库开源情况统计

2023 年，中国国产数据库开源趋势进一步加快。根据墨天轮国产数据库排行榜数据显示，截至 2023 年 12 月底，已有 51 款国产数据库产品开源，具体的国产数据库开源情况如下表所示。

表 6 2023 年国产数据库开源情况统计表

序号	名称	贡献者	开源地址	Star 数量	开源时间
1	TiDB	PingCAP	github.com/pingcap/tidb	35.6k	2015 年
2	Milvus	Zilliz	github.com/milvus-io/milvus	25.1k	2019 年
3	TDengine	涛思数据	github.com/taosdata/TDengine	22.5k	2019 年
4	TiKV	PingCAP	github.com/tikv/tikv	14.1k	2016 年
5	Apache Doris	百度	github.com/apache/incubator-doris	10.6k	2017 年
6	NebulaGraph	悦数科技	github.com/vesoft-inc/nebula	9.8k	2019 年
7	StarRocks	镜舟科技	github.com/StarRocks/starrocks	7.2k	2021 年
8	OceanBase	OceanBase	github.com/oceanbase/oceanbase/	7.1k	2021 年
9	Databend	数变科技	github.com/datafuselabs/databend	6.9k	2021 年
10	Pika	奇虎 360	github.com/Qihoo360/pika	5.5k	2015 年
11	AliSQL	阿里巴巴	github.com/alibaba/AliSQL	4.7k	2016 年
12	IoTDB	清华大学	github.com/apache/iotdb	4.2k	2017 年
13	Apache Kylin	易趣	github.com/apache/kylin	3.6k	2014 年
14	GreptimeDB	格睿云	github.com/GreptimeTeam/greptime-db	3.3k	2022 年

15	PolarDB	阿里云	github.com/ApsaraDB/PolarDB-for-PostgreSQL	2.7k	2021 年
16	Kvrocks	美图	github.com/KvrocksLabs/kvrocks	2.7k	2019 年
17	Tendis	腾讯	github.com/Tencent/Tendis	2.7k	2020 年
18	LinDB	饿了么	github.com/lindb/lindb	2.8k	2019 年
19	HugeGraph	百度	github.com/apache/incubator-hugegraph	2.5k	2018 年
20	TerarkDB	字节跳动	github.com/bytedance/terarkdb	2k	2020 年
21	Vearch	京东	github.com/vearch/vearch	1.8k	2020 年
22	RadonDB	青云	github.com/radondb/radon	1.8k	2018 年
23	MatrixOne	矩阵起源	github.com/matrixorigin/matrixone	1.6k	2021 年
24	OpenMLDB	第四范式	github.com/4paradigm/OpenMLDB	1.5k	2021 年
25	CnosDB	诺司时空	github.com/cnosdb/cnosdb	1.5k	2021 年
26	FlashDB	阿明克	github.com/armink/FlashDB	1.5k	2020 年
27	CovenantSQL	子午星辰	github.com/CovenantSQL/CovenantSQL	1.5k	2018 年
28	ByConity	字节跳动	github.com/ByConity/ByConity	1.4k	2023 年
29	TDSQL	腾讯云	github.com/Tencent/TBase	1.4k	2019 年
30	TensorBase	致大尽微	github.com/tensorbase/tensorbase	1.4k	2021 年
31	openGauss	华为	gitee.com/opengauss/openGauss-server	1.3k	2020 年
32	BaikalDB	百度	github.com/baidu/BaikalDB/	1.1k	2018 年
33	云树®Shard	爱可生	github.com/actiontech/dble	1.1k	2016 年
34	openGemini	华为	github.com/openGemini/openGemini	958	2022 年
35	Apache HAWQ	偶数科技	github.com/apache/hawq	696	2015 年
36	StoneDB	石原子科技	github.com/stoneatom/stonedb	841	2022 年
37	gStore	北京大学	github.com/pkumod/gStore	708	2014 年
38	IvorySQL	瀚高	github.com/IvorySQL/IvorySQL	672	2021 年
39	SequoiaDB	巨杉数据库	github.com/SequoiaDB/SequoiaDB	311	2015 年
40	DingoDB	九章云极	github.com/dingodb/dingo	273	2021 年
41	TenDB Cluster	腾讯	github.com/Tencent/TenDBCluster-TSpider	134	2020 年
42	AwaDB	阿哇科技	github.com/awa-ai/awadb	146	2023 年

43	Claims	华东师范大学	github.com/dase/CLAIMS	115	2016 年
44	PinusDB	巨松软件	github.com/pinusdb/pinusdb	109	2018 年
45	CloudberryDB	酷克数据	github.com/cloudberrydb/cloudberrydb	104	2023 年
46	GreatSQL	万里数据库	github.com/GreatSQL/GreatSQL	106	2021 年
47	CeresDB	蚂蚁集团	gitee.com/ceres-db/ceresdb	71	2022 年
48	Palo	百度	github.com/baidu/palo	52	2017 年
49	Yukon	超图软件	gitee.com/opengauss/yukon	43	2022 年
50	He3DB	移动云	gitee.com/he3db/	62	2022 年
51	OushuDB	偶数科技	github.com/oushu-io/flink2oushodb	4	2018 年

二、Github 上最火的国产开源数据库：TiDB

2023 年 PingCAP TiDB 的 Github Stars 数突破 35.6k 万,连续多年蝉联 GitHub 上热度最高的中国开源数据库。TiDB 也是唯一入选 DB-Engines 全球数据库排行榜 TOP100 的国产开源数据库。

三、2023 年 Github Stars 数增长最多的国产开源数据库 TOP3

- Zilliz 旗下向量数据库 **Milvus** 是 2023 年 Github Stars 数增长最多的国产开源数据库，2023 年 Github Stars 数增长 10.5k，Github Stars 数高达到 25.1k，仅次于 PingCAP 旗下的 TiDB。
- **Apache Doris** 是 Github Stars 数增长第二的国产开源数据库，2023 年 Github Stars 数增长 3.9k，Github Stars 数高达到 10.6k，是 Github 上排名第 5 的国产开源数据库。
- Linux 基金会项目 **StarRocks** 是 Github Stars 数增长第三的国产开源数据库，2023 年 Github Stars 数增长 3.9k，Github Stars 数高达到 7.2k，是 Github 上排名第 7 的国产开源数据库。

四、2023 年开源的国产开源数据库

2023 年 5 月，字节跳动将 ByteHouse 内核向社区开源为 **ByConity**，并于近日正式官宣布 0.1.0 版本。ByConity 定位为开源的云原生数据仓库，采用 Apache 2.0 许可协议，基于 ClickHouse 内核，但采用了存储计算分离的全新架构，支持多个关键功能特性，如存储计算分离、弹性扩缩容、租户资源隔离和数据读写的强一致性等。

通过利用主流的 OLAP 引擎优化，如列存储、向量化执行、MPP 执行、查询优化等，ByConity 可以提供优异的读写性能。ByConity 适合用于 Online Analytical Processing (OLAP) 场景和轻载数仓场景，包括但不限于交互式分析、实时日志监控、流数据处理和分析等。



北京阿哇科技有限公司 (Awa AI) 成立于 2023 年 4 月，并于 5 月开源了 **AwaDB**。AwaDB 是一款面向大模型 /AIGC 场景的 AI Native 新型向量数据库，基于京东 Vearch 核心向量引擎创新研发。AwaDB 结合大模型 /AIGC 落地场景需求，在用户体验上做了较大的改进，能让用户迅速地在 ChatGPT 等场景中得以应用，目前已经与 langchain, llama_index 等流行的大模型应用框架集成，与大模型及其生态无缝衔接，开箱即用。在用户体验的提升改进过程中，对向量引擎内核也做了一些重大的改进及增强，比如增加了文本过滤功能，标量字段行存变列存等等。



2023 年 7 月 14 日，**Cloudberry Database v1.0.0** 发布。CloudberryDB 是酷克数据面向分析和 AI 场景打造的下一代统一型开源数据库，搭载了 PostgreSQL 14.4 内核，兼容 PostgreSQL 和 Greenplum 生态，采用 Apache License 2.0 许可协议。CloudberryDB 支持丰富的数据类型和数仓/AI 混合负载，可开展 SQL 分析、机器学习、全文检索、HTAP 等任务，通过数据存储加密、联合身份验证等技术手段，帮助企业更方便地自建高效稳定的数据底座。



五、入选“2022-2023 年度世界开源贡献榜”的国产数据库

2023 年 11 月 7 日，Bench Council（国际测试委员会）公布了“2022-2023 年度世界开源贡献榜”，该榜单号称“只以贡献分高下”。中国的国产图数据库 gStore、时序数据库 CeresDB、GreptimeDB 与 流数据库 RisingWave 入选“2022-2023 年“开源系统杰出成果”；北京大学邹磊作为 gStore 主要学术贡献者入选“开源系统杰出人才”。

六、开源数据库面临诸多风险

开源的好处众所周知，在此不再赘述。但事物有两面性，我们在享受开源带来收益的同时，也需要警惕开源引入的风险。清楚的认识到的开源风险并加以防范，才能让开源走的更远，行业发展更加健康。

（一）知识产权与开源协议

开源软件一般分为自主开源与被动开源。国内数据库厂商完全自主研发的数据库，采用开源模式进行推广的是主动开源；国内数据库厂商基于国外开源数据库内核二次开发，因为需要遵守国外开源数据库开源协议必须将二次开发的源代码公开的，是被动开源。虽然这两种数据库都统称为开源数据库，但在掌握核心技术和知识产权方面存在巨大差异。自研数据库厂商主动开源是授权方掌握主导权，具备完整的知识产权；被动开源数据库厂商则会被国外开源协议限制是被授权方，存在知识产权侵权风险。

开源软件的知识产权通常通过开源协议进行限制，开源协议本质是开源软件的公共授权合同。据不完全统计全球开源协议有 2000 多种，不同的开源协议的要求和限制差别较大。开源软件可能引用其他开源产品的部分或全部代码，因为这些开源软件遵循不同的开源协议，这种相互嵌套会导致某款开源数据库法律声明中看似使用单一宽松开源协议，实际上内含多种限制性开源协议，不同开源许可协议之间混合或者不兼容使得终端用户难以取得完整的开源数据库合法使用授权。

另外，实际控制开源项目的开源基金会或者商业机构可以根据其商业意图更换开源协议，例如国外知名数据库厂商 MongoDB 曾将 AGPL 协议变更为更为严格的 SSPL 协议，这种开源协议变更甚至开源转闭源为终端用户选型增加了难度。

违反开源协议容易造成违约行为，导致数据库厂商和终端用户面临知识产权侵权诉讼风险，可能会被权利所有人提起专利诉讼或收取授权费用，影响较大的案例会导致商誉受损。

（二）安全漏洞与断供停服

开源软件因为代码公开可见，其设计缺陷或错误代码导致的安全漏洞极易被黑客发现并利用，毫无秘密可保。如果终端用户在使用开源数据库时未及时发现和修复安全漏洞，将面临数据泄露等重大安全威胁。大部分终端客户内部开源治理制度及管控流程还不够成熟，缺乏配套的专业人员和工具，在关键业务系统上使用开源数据库容易被植入木马和发生漏洞攻击事件，无法从根本上保障数据安全。

当前国际形势异常复杂，开源软件事实上已有国界，跨国的开源软件供应链随时可以被切断，断供是制约开源软件技术交流和发展的一个显性风险。

开源数据库厂商的开发与维护投入非常巨大，开源数据库产品停止开发和维护是正常商业现象。以当前比较流行的美国甲骨文公司 MySQL 开源数据库为例，2023 年 10 月 MySQL5.7 稳定版本停服在国内引起广泛关注，国内终端用户被迫将其生产环境 MySQL 数据库升级至更高版本或者进行系统迁移。一旦开源数据库某个版本或者整条产品线停服，终端用户将无法得到相关的技术支持，可能因为缺乏维护导致安全漏洞无法及时修复，也会严重影响存量信息系统的运行维护，带来大量的系统更新需求和额外的成本及风险。

8、数据行业分析报告

2023 年，中国数据库行业在技术创新和市场认可方面取得了显著成就。本土厂商如 OceanBase、PingCAP 和腾讯云的数据库产品因卓越的安全性能和客户体验获得国际权威机构如 IDC 和 Gartner 的认可。同时，随着湖仓一体架构和图数据库技术的快速发展，中国数据库市场规模持续扩大，尤其在金融和云服务领域表现突出。腾讯云数据库 TDSQL 在国内外市场都取得了显著的成绩，阿里云更是在全球数据库市场份额中跻身前十，显示了中国数据库行业的国际竞争力。

以下是摘选的部分数据库行业分析报告的内容：

一、市场份额报告

- **Gartner Market Share Analysis: Database Management Systems, Worldwide, 2022**

2023 年 7 月 6 日消息，Gartner 发布的《Market Share Analysis: Database Management Systems, Worldwide, 2022》（2023 年 6 月）报告显示，“**2022 年数据库管理系统市场的总体增长率为 14.4%，超过了整体软件市场 11.3% 的增长。**”，**腾讯在全球市场中收入增长率达 41%，位列国内所有数据库厂商之首。**

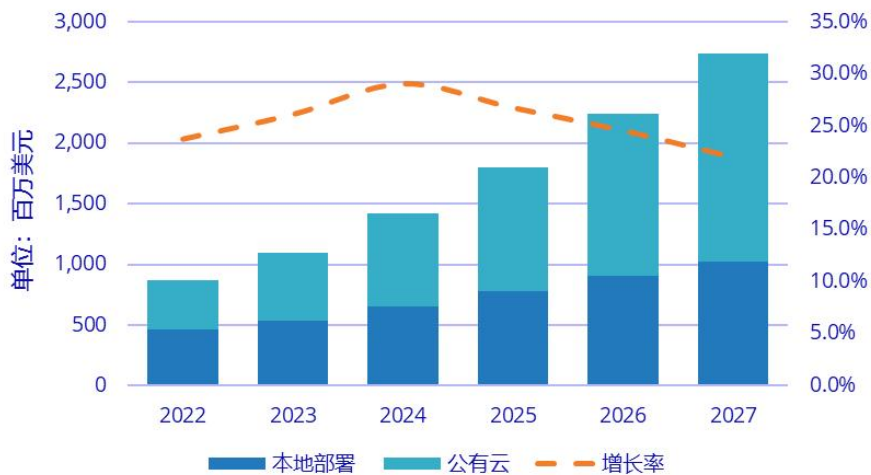
- **IDC 2022 年下半年中国数据仓库软件市场跟踪报告**

8 月 18 日，IDC《2022 年下半年中国数据仓库软件市场跟踪报告》显示，**2022 年中国数据仓库软件市场规模为 8.7 亿美元，同比增长 23.7%**。其中，本地部署数据仓库软件规模为 4.6 亿美元，同比增长 12.5%；公有云数据仓库软件规模为 4.1 亿美元，同比增长 39.3%。**2022 年下半年，华为云 GaussDB(DWS)在中国数据仓库软件整体市场和本地部署排名“双第一”，GaussDB(DWS)已连续两年蝉联本地部署模式市场份额第一，并在 2022 年 H2 整体市场排名中位居首位，占比高达 19.6%。**

IDC 预测，到 2027 年，中国数据仓库软件市场规模将达到 27.3 亿美元，2022-2027 的 5 年市场年复合增长率（CAGR）为 25.7%。



中国数据仓库软件市场规模预测, 2022H2

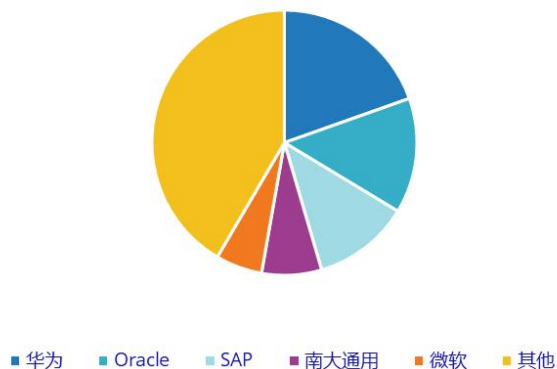


拓展: 从部署模式来看, 出于数据安全和合规性的考虑, **本地部署模式的数据仓库产品**仍将是政府、金融、能源, 以及大型企业的首选。同时, 传统国际数仓品牌由于支撑当前海量数据分析场景的成本较高, 企业正在寻求替代方案, 而本土厂商的性价比更具优势。

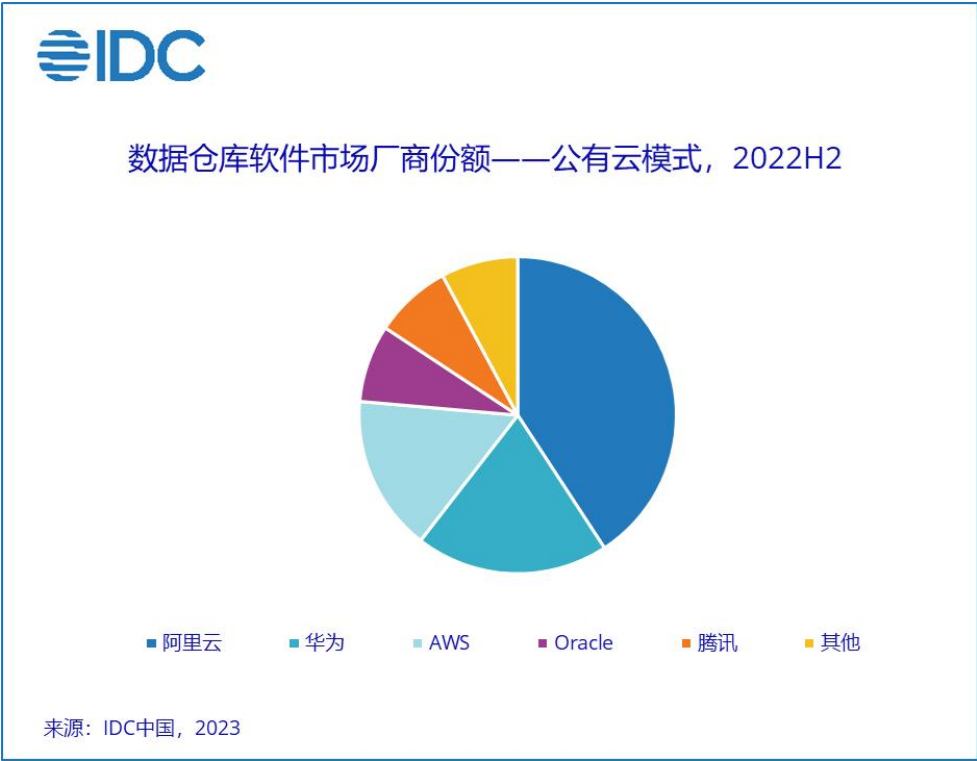
2022 下半年中国数据仓库本地部署模式市场厂商份额情况如下:



数据仓库软件市场厂商份额—— 本地部署模式, 2022H2



拓展：与本地部署市场相比，公有云数据仓库服务的市场集中度更高，前五名厂商份额共计超过 90%。中国泛互联网行业 and 传统企业的互联网业务，对公有云数据服务有着强烈的需求，相关企业过去对大数据服务、关系型数据库服务已经有广泛的采用，并在公有云上积累了大量的数据，为云上数仓的使用创造了前提和基础。预计到 2027 年，公有云数据仓库市场规模占比将从 2022 年的 47.2% 上升到 62.6%。2022 年下半年中国数据仓库公有云模式市场厂商份额情况如下：



● IDC 2023 年上半年中国关系型数据库软件市场跟踪报告

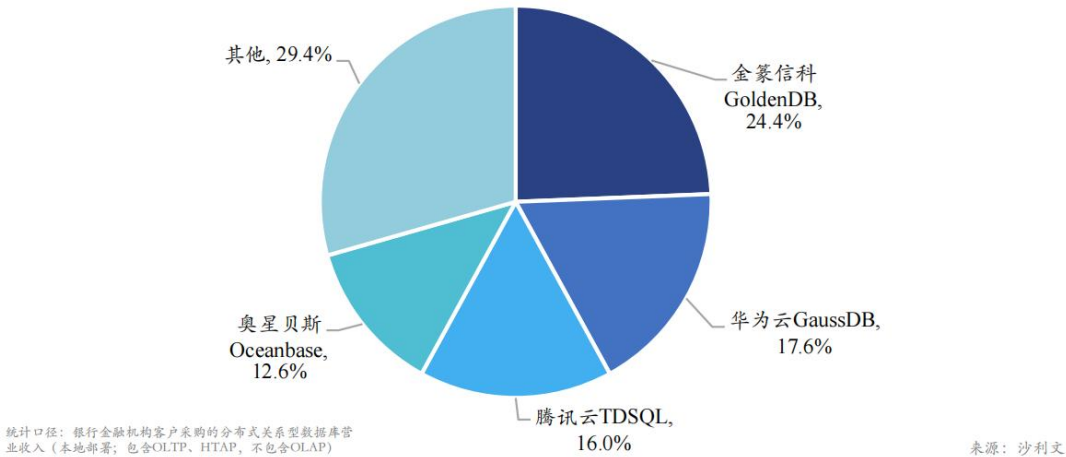
12 月 8 日消息，IDC 发布《2023 年上半年中国关系型数据库软件市场跟踪报告》。报告显示，2023 年上半年中国关系型数据库整体市场规模为 17.5 亿美元，同比增长 13%。选择公有云模式部署数据库已成为主流，该模式同比增长 19%，达 11.3 亿美元，市场规模是本地部署模式的 2 倍左右。2023 年上半年，阿里云以 27% 的整体市场份额（公有云+本地部署模式）位居第一，占比高于第二名和第三名的总和。公有云部署模式中，阿里云占比 39%，以绝对优势连续 9 年蝉联第一。华为云 GaussDB 数据库以 15.47% 的市场份额，位列本地部署模式国产厂商第一。

● 沙利文 - 中国金融数据库行业研究（2023）

2023 年 12 月，沙利文(FROST&SULLIVAN)发布《中国金融数据库行业研究（2023）》，报告显示，2022 年中国金融级分布式数据库市场份额前四名分别是 GoldenDB（25%）、华为云 GaussDB（14.9%）、OceanBase（13.0%）、腾讯

云 TDSQL（9.6%），前四名同占据近 60% 的市场份额。沙利文预测金融级分布式数据库市场规模将从 2023 年的 16.5 亿元，在 2027 年增长至 50.8 亿元，复合年增长率达 32.6%。

2022年中国银行业金融级分布式数据库市场份额（以营收计）



拓展：其中金篆信科 GoldenDB 在金融领域已形成领先落地深度与广度，银行业金融级分布式数据库市场份额排名第一，占比 24.4%；银行核心系统（不包括次核心）市场投产数量行业占比 50%，次核心及非银核心系统市场投产数量占比 32%，均为行业第一！这是 GoldenDB 继过去两年蝉联金融级分布式数据库第一之后，再次获得沙利文的认可。

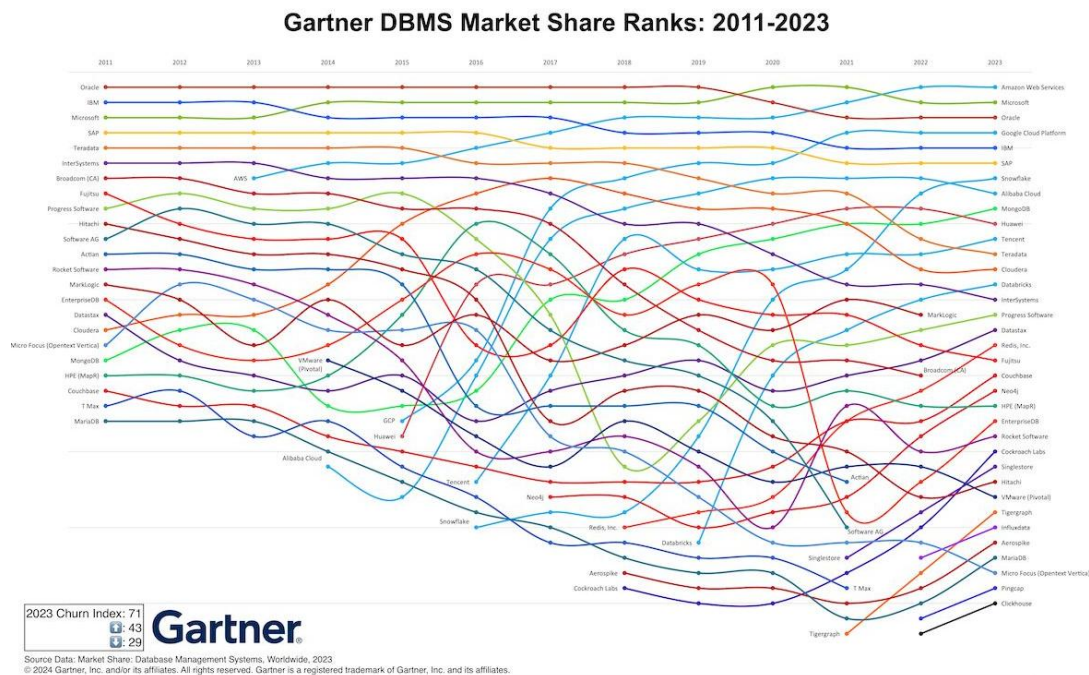
● 沙利文 - 重点行业数据库应用调研报告

12月28日，在 openGauss 2023 峰会上，沙利文杨晓骋发布《重点行业数据库应用调研报告》，报告指出 2023 年中国数据库市场线下集中式 openGauss 系新增市场份额达 21.9%，已规模应用于金融、政府、电信、能源、制造、公路水运、邮政、教育等十大关键行业核心场景。这标志着 openGauss 已跨越生态拐点，正式踏入生态发展期。



- **Gartner 2023 年数据库市场份额报告**

2024 年 4 月 25 日消息，Gartner 发布了“2023 年数据库市场份额报告”，前五名分别是亚马逊、微软、Oracle、谷歌、IBM，排名相比去年没有变化。中国今年有 4 家厂商上榜：阿里云下降一位排名第八，华为下降两位排名第十，腾讯上升一位排名第 11，PingCAP 新上榜，排名第 34 名。



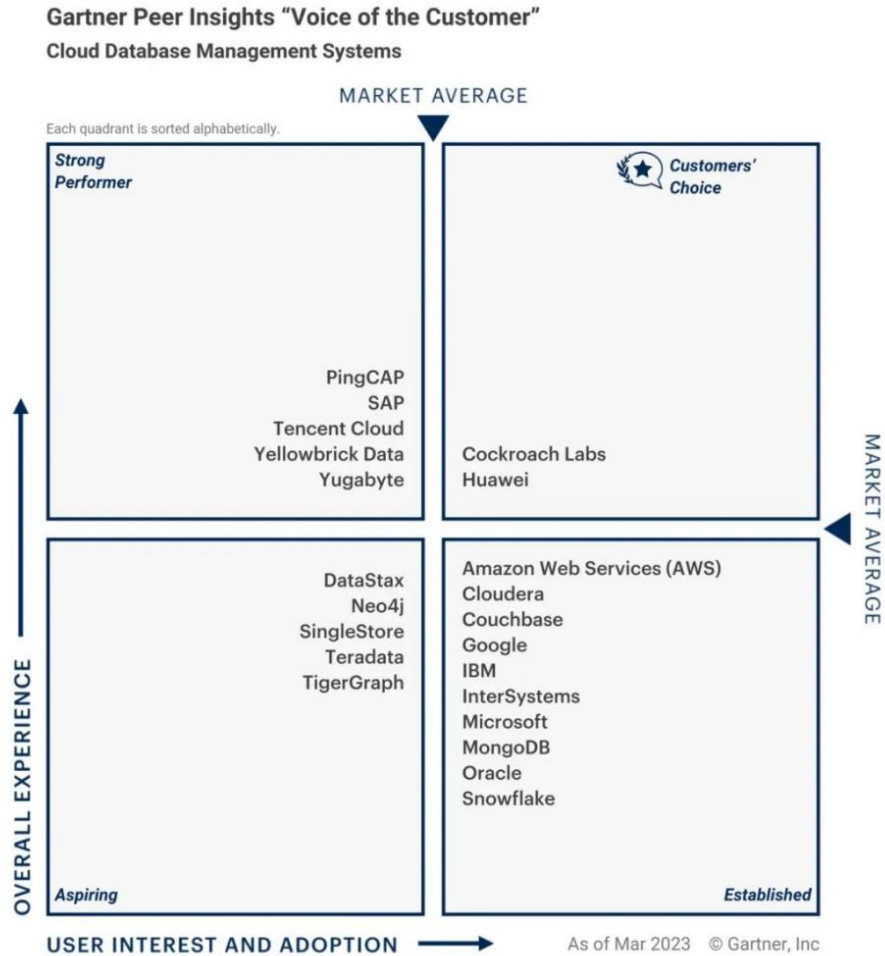
二、影响力报告

- **IDC 中国数据库原生安全能力洞察 2022**

5 月 10 日消息，近日，国际权威市场研究机构 IDC 正式发布《中国数据库原生安全能力洞察 2022》研究报告，OceanBase 凭借领先的数据一致性、数据访问控制、数据加密、高可用等数据库原生安全能力，作为中国数据库的代表厂商入选此报告。入选该报告，意味着 OceanBase 原生安全能力的卓越表现得到了国际权威机构的认可。

- **Gartner Peer Insights 云数据库管理系统客户之声**

6 月 5 日消息，近日，Gartner 发布“Gartner Peer Insights™ 云数据库管理系统客户之声”报告，PingCAP 被评为“卓越表现者”。连续两年获评“卓越表现者”最高分，表明 TiDB 在整体体验方面超过了市场平均水平。腾讯分布式 SQL (TDSQL) 被列为“卓越表现者”，97% 的客户表示愿意推荐腾讯云的数据库产品和服务，被评为亚太客户之选。



拓展：本次报告展示了来自全球的金融、服务、游戏、物流、零售等多个行业的客户在产品功能、部署运维、服务支持等多个维度上对 TiDB 产品及服务的综合评估结果。这些评价充分体现了全球客户对 TiDB 的高度认可，97% 的客户愿意推荐同行业客户尝试使用 TiDB，使得 TiDB 成为名副其实的“卓越表现者”（Strong Performer）。腾讯云数据库管理系统（DBMS）产品矩阵包括分布式 OLTP 数据库腾讯分布式 SQL（TDSQL）、KV 存储 KeeWiDB、时序 CTSD、图 KonisGraph、腾讯大数据套件 TBDS。

● 中国信通院 - 数据库发展研究报告（2023 年）

7 月 4 日，“2023 可信数据库发展大会”在京召开，中国信通院正式发布《数据库发展研究报告（2023 年）》暨《中国数据库产业图谱（2023 年）》。《报告》涵盖中国数据库产业发展态势介绍、数据库政策解读、数据库技术趋势剖析等内容；《图谱》则全面客观地展现了中国数据库产业中的关键领域、环节和代表企业，为社会各​​界提供了中国数据库发展的一站式信息库。

● IDC Technology Assessment: 湖仓一体数据平台技术能力评估报告，2023

7月13日消息，IDC发布了《IDC Technology Assessment: 湖仓一体数据平台技术能力评估报告，2023》，从数据管理、数据存储、数据开发、数据安全、工程化落地、生态、覆盖行业等方向进行评估，并筛选出10家代表厂商，分别是阿里云、柏睿数据、滴普科技、华为云、金山云、巨杉数据库、科杰科技、网易数帆、星环科技、亚信科技（按照首字母顺序排序）。

拓展：IDC 调研显示，有 66.9% 的企业了解湖仓一体架构，有 85% 的企业正在部署或考虑评估部署湖仓一体架构。IDC 数据统计，2022 年中国大数据市场总体 IT 投资规模约 170 亿美元，并在 2026 年增至 364.9 亿美元，实现规模翻倍，与全球总规模相比，中国市场在五年预测期内占比持续增高，有望在 2024 年超越亚太（除中日）总和，并在 2026 年接近全球总规模的 8%。湖仓一体作为核心数据管理架构，是后疫情时代企业布局和采购的重点，将对海量多模数据管理和价值挖掘产生重要影响。

● **IDC MarketScape: 中国图数据库市场 2023 年厂商评估**

7月15日消息，IDC发布了《IDC MarketScape: 中国图数据库市场 2023 年厂商评估》市场研究报告，报告显示，蚂蚁集团自研的企业级图数据管理平台 TuGraph 跻身"领导者"象限。



拓展：本次报告 IDC 主要评估了中国市场上 12 家典型的企业级图数据库厂商（创邻科技、Fabarta、华为云、杭州悦数、海致科技、蚂蚁集团、明略科技、Neo4j、蜀天梦图、TigerGraph、星环科技、赢图等），类型覆盖互联网厂商、云服务厂商、大数据厂商等。

根据 IDC 数据，2022 年中国图数据库市场规模达到 2.4 亿元人民币，预计 2023 年整体市场规模将超过 4 亿元人民币，呈现快速增长的态势。从整体市场来看，市场拓展至多行业需求，金融仍是主要客户渠道，数字化转型是当前市场化主要推手。

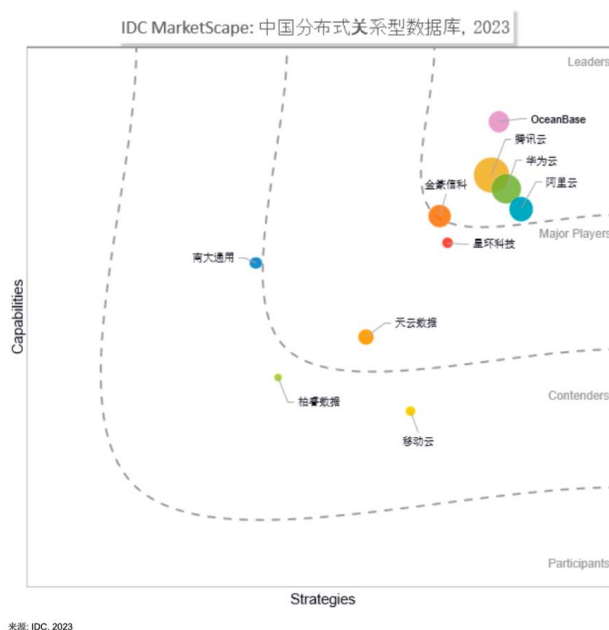
- **Gartner 2023 中国数据库管理系统市场指南**

8月9日消息,日前,Gartner 发布《2023 中国数据库管理系统市场指南》(Market Guide for DBMS, China ,2023), 在中国范围内评估并重点推荐了 36 家极具实力的企业, 柏睿全内存分布式数据库 RapidsDB、DolphinDB、亚信 AISWare AntDB 数据库入选典型产品。

拓展: Gartner 在该指南中认为, 随着本地云计算多样化市场格局的形成, 以及技术自给自足趋势的驱动, 中国数据库管理系统 (DBMS) 市场正在快速增长与变化。Gartner 数据显示, 2022 年, 中国 DBMS 市场规模增长至 63.24 亿美元, 比 2021 年增长了 11.4%; 其中, 中国供应商的市场份额增至 50.1%, 首次超过 50%。并预测到 2025 年, 海外供应商在中国分析型数据库的市场份额将仅占 30%, 在交易型数据库市场份额约为 50%。

- **IDC MarketScape: 中国分布式关系型数据库 2023 年厂商评估**

12月6日消息, IDC 近日发布的《IDC MarketScape: 中国分布式关系型数据库 2023 年厂商评估》报告显示, OceanBase、腾讯云、华为云、阿里云、金篆信科等位居中国分布式关系型数据库“领导者”类别。



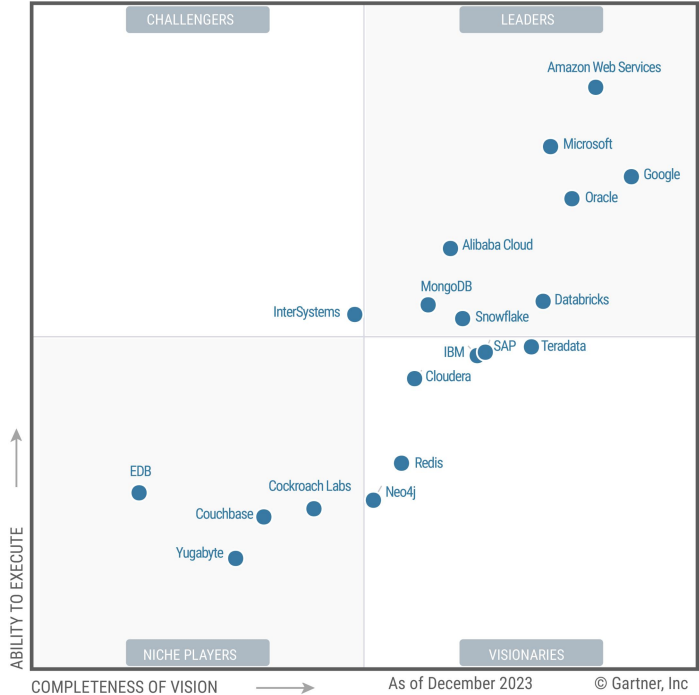
- **沙利文&头豹研究院 - 2023 年中国金融级分布式数据库市场报告**

12月12日消息, 沙利文联合头豹研究院发布《2023 年中国金融级分布式数据库市场报告》。由“创新指数”和“增长指数”综合评分, 华为云 GuassDB、金篆信科 GoldenDB、腾讯云 TDSQL、奥星贝斯 OceanBase、PingCAP TiDB 位列中国金融级分布式数据库市场领导者梯队。

● Gartner 2023 年云数据库管理系统魔力象限

12 月 18 日，Gartner 发布了“2023 年云数据库管理系统魔力象限”，阿里云入选“领导者”象限，是该魔力象限唯一入选的中国厂商，其数据库管理系统产品包括 RDS 和 PolarDB。用于分析场景的产品有 AnalyticDB 和 MaxCompute，此外，阿里云还提供 Lindorm、GDB 和 Tair，用于非关系型场景。在继去年华为云退出魔力象限后，今年腾讯云也退出了该魔力象限。“象限”报告最后的“Honorable Mentions”部分，提到的中国厂商包括 OceanBase、华为云、腾讯云、PingCAP（TiDB）。

Figure 1: Magic Quadrant for Cloud Database Management Systems



Source: Gartner

拓展：2023 Gartner 云数据库管理系统魔力象限入选名单如下：

- 领导者：Amazon Web Services、Microsoft、Google、Oracle、Alibaba Cloud、MongoDB、Databricks
- 挑战者：InterSystems
- 远见者：Teradata、SAP、IBM、Cloudera、Redis、Neo4j
- 小众玩家：Cockroach Labs、Couchbase、Yugabyte、EDB

总结：

2023 年，中国数据库行业实现了显著的技术创新和市场拓展，国际竞争力显著提升。国内领先厂商如 OceanBase、PingCAP、腾讯云和华为云等不仅巩固了国内市场的领导地位，还在全球范围内获得了广泛认可。在金融级分布式数据库这一细分市场，

GoldenDB、GaussDB、OceanBase 和 TDSQL 等产品以其卓越的性能和安全性，确立了行业领导地位。

展望未来，预计中国数据库市场将维持强劲的增长势头。随着数字化转型的深入，本土厂商有望进一步巩固并扩大其在全球数据库产业中的影响力。特别是在金融、云服务和大数据等关键领域，中国数据库厂商预计将发挥更加重要的作用，引领整个行业向更高效、更安全、更智能的方向发展，为全球数据管理技术的进步贡献中国力量。

9、数据库行业大事记

一、大事记——上市进程

随着我国科技实力的不断提升，国产数据库公司逐渐崭露头角，加速上市进程。2023 年，航天软件、第四范式、武汉达梦数据，以及柏睿数据等知名数据库公司纷纷迈出了重要步伐。

- **航天软件上交所成功上市，年营收 19 亿，市值 114 亿**

2023 年 5 月 24 日，北京神舟航天软件技术股份有限公司在上海证券交易所科创板上市。航天软件本次发行股票 1 亿股，每股发行价格 12.68 元，发行市盈率 133.70 倍，募集资金总额 12.68 亿元，超募 7.17 亿元，当日开盘价 25.19 元，涨幅 98.66%，充分体现了资本市场对国产数据库、工业软件和航天软件市场价值的认可。航天软件是我国四大国产关系型数据库厂商之一，其神通数据库自 1980 年开始研发攻关，先后受到 5 项国家“核高基”重大专项支持（3 项独立或牵头），参与制定 5 项国家标准，是国家重点支持的数据库产品研发国家队。

- **第四范式港交所上市，市值 230 亿港元**

2023 年 9 月 28 日，人工智能公司第四范式正式登陆港交所上市，开盘价报 63.1 港元/股，高开 13.49%。招股书显示，2020 年至 2022 年，公司营收快速增长，分别为 9.422 亿元、20.184 亿元、30.826 亿元，同比增长率分别为 114.2%、52.7%。2023 年一季度的营收也同比增长了 33.6%。营收增长的一大原因在于不断扩大的用户群。虽然报告期内公司尚未实现盈利，但其 2022 年亏损幅度已有所收窄。

- **柏睿数据上市辅导备案获北京证监局受理登记**

2023 年 10 月 17 日，柏睿数据在北京证监局办理辅导备案登记，拟首次公开发行股票并上市。这家以数据库为核心的“Data+AI”数据智能基础软件公司在融资历程中，获得了多家知名投资机构的青睐，累计融资额高达 4 亿元以上。2022 年，公司完

成了过亿元 D 轮融资，由中科海创、上海国际、北科建以及朝科创共同出资完成。此次辅导备案得到北京证监局受理登记，意味着该公司在资本市场的道路上也取得了重要进展。

- **达梦数据库科创板上市已获证监会批准**

2023 年 12 月 20 日，武汉达梦数据库股份有限公司科创板上市已获证监会批准。此次 IPO，达梦数据计划募资 23.51 亿元，拟公开发行股票数量为不超过 1900.00 万股，且不低于发行后公司总股本的 25%。此次募集的资金将用于达梦集群数据库管理系统、高性能分布式关系数据库管理系统、新一代云数据库产品的优化升级，以及达梦中国数据库产业基地和达梦研究院建设等项目。作为我国数据库领域的佼佼者，武汉达梦一直致力于自主研发与创新。在上市过程中，该公司得到了监管机构和市场各方的认可，预示着其未来发展前景广阔。

整体来看，这几家国产数据库公司在上市进程中均表现出良好的发展势头。随着政策的扶持和市场需求的不断增长，我国数据库产业将持续发展，这些公司有望在竞争中脱颖而出，为我国科技产业的繁荣做出更大贡献。同时，上市将为这些企业提供更多的资金和资源支持，助力它们加大研发投入、提升技术实力、扩大市场份额。在未来的市场竞争中，国产数据库公司有望占据一席之地，为全球市场带来新的活力。

二、大事记——国际视野

- **IBM 收购 GraphQL server 开发商 StepZen**

2023 年 2 月 13 日消息，IBM 宣布已完成对 StepZen 公司的收购，这家公司开发了一个具有独特架构的 GraphQL 服务器，可以帮助开发人员使用更少的代码来快速构建 GraphQL API。这是 IBM 步入 2023 年的首项收购，自 Arvind Krishna 于 2020 年 4 月担任首席执行官以来，IBM 已收购了 30 家公司，以强化其混合云与人工智能（AI）的能力。

- **InfluxDB 厂商完成 5100 万美元 E 轮融资**

2023 年 2 月 14 日消息，时序数据库 InfluxDB 厂商 InfluxData 完成 5100 万美 E 轮融资，融资规模扩大至 1.71 亿美元。InfluxDB 一直在 DB-Engines 时序数据库排行榜上排名第一，是一款非常受欢迎的时序数据库，适合存储设备性能、日志、物联网传感器等带时间戳的数据。

- **数仓巨头 Teradata 退出中国市场**

2023 年 2 月 15 日，大数据分析/数仓软件巨头 Teradata 宣布退出中国市场，后续将进入中国公司的关闭程序！Teradata 天睿公司（纽交所代码：TDC），是美国前

十大上市软件公司之一。经过逾 30 年的发展，Teradata 天睿公司已经成为全球最大的专注于大数据分析、数据仓库和整合营销管理解决方案的供应商。2022 年，Teradata 全年营收 17.95 亿美元，同比减少 6.4%；其中亚太及日本市场贡献营收 2.92 亿美元，同比减少 11.5%。

- **DragonflyDB 获得 2100 万美元融资并发布 1.0 版本**

2023 年 3 月 21 日，内存数据库初创公司 DragonflyDB Inc 宣布获得 2100 万美元融资，同时推出其数据库的最新版本 DragonflyDB 1.0，该版本增加了几个新的可靠性和数据管理功能。DragonflyDB 将使用新的资金来增强其数据库。DragonflyDB 计划增加对更多类型工作负载的支持并简化用户体验。此外，其打算推出一个托管版本，使组织更容易使用数据库。

- **四家向量数据库公司共获得 10 多亿融资**

2023 年 4 月，为 ChatGPT 等生成式 AI 应用提供向量搜索、向量数据存储、向量嵌入等功能的向量数据库赛道突然走红，四家公司 Qdrant、Pinecone、Chroma 和 Weaviate 最近共获得 10 多亿融资。4 月 10 日，Chroma 在 Quiet Capital 的 Astasia Myers 的带领下筹集了 1800 万美元的种子轮。4 月 19 日，开源向量数据库初创公司 Qdrant 宣布从主要投资者 Unusual Ventures 获得 750 万美元的种子融资。4 月 22 日，向量数据库平台（vector database）Weaviate 宣布获得 5000 万美元（约 3.5 亿元）B 轮融资。4 月 28 日，向量数据库平台 Pinecone 宣布获得 1 亿美元（约 7 亿元）B 轮融资。

- **Oracle Database 23c 免费开发者版正式发布**

2023 年 4 月 3 日，甲骨文宣布推出免费版 Oracle Database 23c。Oracle Database 23c 免费版 — 开发者版本满足了全球开发人员和组织日益增长的访问 Oracle Database 23c “App Simple” 中最新特性的需求。Oracle Database 23c 免费版 — 开发者版本以 Docker 镜像、VirtualBox 虚拟机和 Linux RPM 安装文件的形式提供，包括 JSON 关系二元性、JavaScript 存储过程、JSON 模式、操作性属性图、Oracle Kafka API、SQL 域、注释等七大突破性的新功能。

- **阿里云与 MongoDB 续签战略合作协议**

2023 年 5 月，阿里云宣布与开发者数据平台公司 MongoDB 续签战略合作协议，双方致力于将 MongoDB 的最新产品成果与阿里云企业级服务相结合，携手深耕广袤的国内市场。阿里云是中国唯一能提供 MongoDB 最新版本的云厂商。从 2021 年 5 月国内独家首发 5.0 版本，到如今的 7.0 版本，阿里云 MongoDB 始终跟随

MongoDB 公司的发布节奏，致力于为开发者带来最新版本的云服务体验。阿里云与 MongoDB 战略合作四周年，目前已经成为 MongoDB 在中国最大的云服务提供商，业务拓展至互联网、游戏、汽车、制造、零售等行业，累计为数万名客户提供 MongoDB 云服务。

- **Oracle Database 19c 支持 ARM 架构**

2023 年 6 月 28 日，Oracle Database 19c 提供了一种新的平台支持——ARM。无论在云端还是用户本地数据中心，开发者都可以将 Oracle Database 19c 部署在当今流行的 ARM 架构上。ARM 技术有着非常多的优势，比如节能和环保，尤其是被广泛应用于各家大规模云计算中心的 Ampere 处理器。据不完全统计，目前已经有 1800 亿个 Ampere 处理器被应用在智能手机、物联网传感器或者其他智能设备上。

- **Oracle 宣布 MySQL HeatWave Lakehouse 上市**

2023 年 7 月 20 日，Oracle 宣布全面推出 MySQL HeatWave Lakehouse，这是业界首创，使客户能够像查询数据库内的数据一样快速地查询对象存储中的数据。MySQL HeatWave Lakehouse 支持多种对象存储文件格式（例如 CSV、Parquet）以及从其他数据库导出文件，并且可以在同一查询中将对象存储文件数据和 MySQL 数据库事务数据组合在一起。HeatWave 直接查询对象存储文件，无需将数据复制到 MySQL 数据库中。因此，MySQL HeatWave Lakehouse 为查询处理的可扩展性和性能、加载数据的速度、集群配置时间以及对象存储中数据查询的自动化设定了新的标准。

- **数据库初创公司 Neon 获得 4600 万美元融资**

2023 年 8 月 1 日，一家提供 Postgres 数据库无服务器选项的初创公司 Neon 完成了 4600 万美元 B 轮融资。投资方为 Databricks Ventures、Founders Fund、Menlo Ventures、纪源资本、Snowflake Ventures、General Catalyst、科斯拉风险投资。

- **实时索引数据库 Rockset 完成 4400 万美元 B+轮融资**

2023 年 8 月 1 日消息，实时索引数据库厂商 Rockset 完成 4400 万美元 B+轮融资，融资规模扩大至 1.05 亿美元。新一轮融资由 Icon Ventures 领投，新投资者 Glynn Capital、Four Rivers、K5 Global 以及现有投资者 Sequoia 和 Greylock 跟投。Rockset 是一个实时索引数据库，可自动在任何数据（包括结构化，半结构化，地理和时间序列数据）上构建 Converged Index，以进行大规模的高性能搜索和分析。它支持对半结构化数据的实时 SQL 查询，因此开发人员可以灵活地构建功能，以快速搜索，聚合和合并来自任何来源的任何类型的数据。

- **Databricks 在 Pre-IPO 轮融资 5 亿美元**

2023 年 9 月 14 日消息,数据分析和人工智能软件制造商 Databricks 已筹集了价值超过 5 亿美元的第一轮融资,估值达到 430 亿美元本轮融资表现突出,尤其是在融资放缓的情况下,许多处于后期阶段的初创公司的估值被大幅减。Databricks 上次融资是在 2021 年 8 月,融资后估值为 380 亿美元,现在该公司的估值增加了 50 亿美元。本轮融资由其现有投资者.Rowe PriceAssociates 领投, Capital ne Ventures、英伟达、老虎环球、摩根士丹利等十几家投资者参投。

三、大事记——商业发展

- **中国软件对易鲸捷投资 6515 万**

2023 年 2 月 16 日,中国软件关联公司对易鲸捷投资 6515 万元。中国软件在与中国电子、中国物流集团共同出资设立物流科技公司的公告中披露,对贵州易鲸捷信息技术有限公司增资事项,按照各方签署的最后交割日 2022 年 11 月 10 日之协议约定,截至最后交割日公司缴付增资款 1791 万元,中软金投缴付增资款 4724 万元,合计持股比例增至 11.09%股权。

- **StarRocks 项目加入 Linux 基金会开源项目**

2023 年 2 月 20 日消息, CelerData 宣布 StarRocks 项目加入 Linux 基金会开源项目系列。StarRocks 是业内第一个解决实时和批量分析中关键技术挑战的分析数据库,将实时分析与 OLAP 数据库和数据湖分析合并到一个引擎和一个数据管道中,利用列式存储引擎和完全矢量化的运算符来提高 CPU 处理性能。

- **佳都科技战略投资睿帆科技,已成为睿帆科技第二大战略股东**

2023 年 3 月 7 日消息,佳都科技战略投资睿帆科技(雪球数据库),目前已经成为睿帆科技第二大战略股东,持股比例达到 17.1%。从 2019 年开始,睿帆科技联合佳都科技,与广州地铁进行合作:借助睿帆自主研发的大数据科学平台 Baymax 及内置的 BaymaxBI,大数据基础组件 Hadoop、Spark、Flink、ES 和 MPP 雪球数据库等,共同搭建地铁数据中台。

- **瑶池数据库打造“云原生+一站式”的数据管理与服务**

2023 年 3 月 24 日,在北京召开的阿里云瑶池数据库峰会上,瑶池首次将云原生数据库 PolarDB 和云原生数据仓库 AnalyticDB 打通融合,形成“云原生一体化”的 HTAP 解决方案。此外,阿里云推出了全新云原生多模数据库 Lindorm AI 引擎。SelectDB 与阿里云正式战略签约,发布“阿里云 SelectDB 版本”。

- **3436 万！腾讯云、PingCAP、中兴联合中标中国建行数据库大单**

2023 年 4 月 24 日，中国建设银行股份有限公司发布《国产数据库-小机下移竞争部分》采购结果公示，腾讯云计算(北京)有限责任公司,平凯星辰(北京)科技有限公司,中兴通讯股份有限公司中标，中标金额合计 3436 万元。

- **紫光股份收购新华三（H3C）**

2023 年 5 月 26 日，紫光股份披露重大资产购买预案和定增预案。公司拟由全资子公司紫光国际信息技术有限公司（“紫光国际”）以现金方式收购新华三集团有限公司（“新华三”）剩余的 49% 的股权，作价 35 亿美元。本次交易完成后，紫光国际将持有新华三 100% 的股权。9 月 24 日晚，紫光股份公告披露，“终止收购新华三剩余股权的相关事项，改为先完成定增股份，再继续推进资产重组”。

- **SelectDB 厂商「飞轮科技」完成新一轮数亿元融资**

2023 年 6 月 1 日，新一代实时数据仓库公司「飞轮科技」宣布完成新一轮融资，融资金额达数亿元。飞轮科技是一家基于 Apache Doris 的商业化公司，为 Apache Doris 用户提供技术支持商业服务，并推出了推出基于 Apache Doris 内核的 SelectDB 商业产品。公司于 2022 年 1 月成立，4 月完成天使轮和天使+轮融资，由 IDG 资本、红杉中国等顶级 VC 投资，融资金额超过 3 亿元。截至目前，飞轮科技的融资总金额已接近 10 亿人民币，创下了近年来开源基础软件领域的新纪录。

- **蚂蚁 TuGraph 工业级流式图计算引擎 TuGraph Analytics 正式开源**

2023 年 6 月 11 日，2023 开放原子全球开源峰会高峰论坛上，蚂蚁 TuGraph 工业级流式图计算引擎 TuGraph Analytics 正式开源。此次开源的工业级流式图计算引擎是蚂蚁从 2017 年开始布局打造，经过五年多工业级应用大考，流式图计算做到了在千亿数据规模的“图”上秒级延迟计算，是蚂蚁风控的核心基础技术，成功解决了金融场景风险分析难、识别率低、时效性差等业界难题。

- **ZILLIZ 旗下开源性能测试工具 Vector DB Bench 正式开源**

2023 年 6 月 20 日，ZILLIZ 旗下开源性能测试工具——Vector DB Bench 正式开源。Vector DB Bench 是为追求高性能数据存储和检索系统的用户设计的开源性能测试工具，它允许用户测试和比较不同向量数据库系统的性能，以确定最适合的数据库系统。Vector DB Bench 目前支持六个向量数据库：Milvus、Zilliz Cloud、Pinecone、WeaviateCloud、QdrantCloud 和 ElasticCloud。

- **阿里云 GraphScope 打破全球最快图计算引擎纪录**

2023 年 7 月 21 日，国际权威图基准测评“LDBC SNB Interactive”榜单更新中出现关键突破：阿里云开源图计算引擎 GraphScope 登顶并打破榜单历史纪录，其单节点执行图数据库查询的吞吐率超过 30000 QPS，性能达此前纪录保持者 2 倍。阿里云 GraphScope 具有一站式、开发便捷、性能极致的特点，已于 2020 年开源。它在业界首次支持 Gremlin 分布式编译优化，提供了各类高效图分析算法和图学习算法，并支持用户自定义算法的即插即用。在保持简单易用的同时，GraphScope 达到了企业级场景的极致性能。

- **四维纵横 YMatrix 完成 B 轮融资**

2023 年 7 月 13 日消息，北京四维纵横数据技术有限公司完成 B 轮融资，投资方为用友网络、龙辕资本。同时，注册资本从 170.5569 万增至 2155.5556 万元，股东新增广州同创壹号股权投资基金合伙企业、用友网络科技股份有限公司（持股 3.09%）、广州龙辕同创贰号股权投资基金合伙企业。

- **Serverless 与 AI 驱动，阿里云瑶池数据库核心能力全面升级**

2023 年 10 月 31 日，在 2023 云栖大会上，阿里云瑶池数据库宣布已全面实现 Serverless 化，并接入通义等大模型能力，大幅提升数据库一站式及智能化水平。同时，PolarDB Always On 系列推出 3 大重磅升级，首个数据智能助手 DMS Copilot 也惊艳亮相。在 Serverless 与 AI 的驱动下，云原生数据库加速向一站式智能数据平台发展演进。

作为 AIGC 应用的基础设施，以 PolarDB、AnalyticDB、Lindorm、RDS 为核心的阿里云瑶池数据库现已全面拥抱向量检索能力，并与通义等大模型深度集成，为用户提供智能化的一站式数据管理平台，加速业务数智创新。为一站式解决开发者的全部需求，阿里云数据库进行了一体化的大幅升级，包括 HTAP 一体化、DB+Cache 一体化、DB+存储一体化等三大能力的全面提升，并在 OLTP、OLAP、NoSQL 等多业务场景融合落地，产品易用性大大提高，进一步简化开发、管理和运维成本。

- **TiDB Serverless 正式商用**

2023 年 7 月 13 日，开源数据库独角兽企业 PingCAP 举办 2023 年用户峰会，并宣布进行了新的产品升级。公司不仅推出了面向中国企业级用户的“平凯数据库”，还正式开启了更低成本的数据库产品 TiDB Serverless 的商用。TiDB Serverless 是一款深度结合 AI 技术的数据库产品，架构简单、门槛较低。

- **阿里云数据库国际峰会首度在印尼召开，AnalyticDB 向量引擎支持定制 AIGC 应用**

2023 年 7 月 25 日，阿里云在印尼雅加达举行首届瑶池数据库国际峰会，面向海外市场正式升级云原生一站式数据管理与服务平台，并重磅推出一体化 HTAP 数据库解决方案。此外，会上阿里云还发布了内置向量引擎功能的最新版云原生数仓 AnalyticDB，用户仅需 30 分钟即可构建专属生成式 AI 应用。据悉，AnalyticDB 现已用于文本搜索、图片搜索等场景，并实现了准确性的大幅提升。

- **腾讯云向量数据库（Tencent Cloud Vector DB）正式上线公测**

2023 年 8 月 1 日，腾讯云向量数据库（Tencent Cloud Vector DB）正式上线公测。作为一款全托管的自研企业级分布式数据库服务，腾讯云向量数据库专用于存储、检索、分析多维向量数据。该数据库支持多种索引类型和相似度计算方法，单索引支持 10 亿级向量规模，可支持百万级 QPS 及毫秒级查询延迟。

- **华为宣布 CANTIAN 引擎开源，携手共建数据库存储新生态**

2023 年 8 月 25 日，在 2023 华为数据存储用户精英论坛上，华为与海量数据、万里数据库、天翼云、南大通用、优炫、爱可生共同发布 CANTIAN 引擎开源。经过与生态伙伴联合测试，华为 CANTIAN 引擎可帮助数据库性能平均提升 5 倍，写性能最大提升 10 倍，单节点故障 30 秒恢复业务。

- **科蓝软件成立子公司上海科蓝数据科技有限公司**

2023 年 9 月 8 日，科蓝软件成立全资子公司上海科蓝数据科技有限公司，专注于国产数据库的经营。科蓝软件 SUNDB 数据库已广泛应用在银行、保险、证券、电信、政府、军工、能源、交通等国家关键信息基础设施领域，在关键业务系统数据库国产化替代方面具有绝对优势。

- **九章云极 DataCanvas 公司完成 3 亿元 D1 轮融资**

2023 年 10 月 8 日，九章云极 DataCanvas 完成总融资额 3 亿元 D1 轮融资。中国电子集团旗下中电智慧基金、华民投、中国太平旗下太平创新、浙江东方旗下东方嘉富等央企国企旗下投资机构，以及卓源资本等专注人工智能赛道的知名财务投资机构参与本轮融资。九章云极 DataCanvas “数据智能基础软件”的定位，不仅通过 AutoML 自动机器学习、AutoDL 自动深度学习和 ModelOps 提供模型运行的全生命周期，更通过其研发的多模向量数据库 DingoDB 落地 Data-Centric AI，让客户实现 “Data、AI and Analytics on One Platform”。

- **四维纵横完成上亿元 B 轮融资，加速打造极简、极速的超融合数据库**

2023 年 10 月 19 日，新一代超融合数据库厂商“四维纵横”完成超亿元人民币的 B 轮融资，本轮融资由用友、顺义产业基金领投，广州同创基金跟投。本轮融资资金将主要用于核心技术研发与商业生态建设，进一步提升 YMatrix 超融合数据库在全场景下的性能与生态优势。作为一家数据库公司，四维纵横成立于 2020 年，以超融合数据库 YMatrix 为企业提供集“分析、事务、时序”为一体的企业级数据库产品服务。

- **腾讯云 TDSQL 再创里程碑！国产数据库首次全量承载券商核心交易**

2023 年 11 月 3 日，东吴证券发布行业首个全面自主创新核心交易系统，并实现全栈上线、全业务覆盖和全客户迁移。系统采用腾讯云 TDSQL，是行业首次由国产数据库全量承载券商核心交易。腾讯云 TDSQL 与系统开发商顶点软件的协同优化，让东吴证券核心系统在满足自主创新要求的同时，性能更加强劲：整体并发处理能力较之前提升 10 倍；单节点业务并发量提升到至少每秒 10 万笔；交易平均延迟从 10 毫秒提升至 1 毫秒以内，支持千万级交易委托；数据库负载在高峰期保持在 10% 以下，能有效满足未来三到五年的需求。

- **深算院 YashanDB 集中签约，关键行业商用提速**

2023 年 11 月 17 日，深圳计算科学研究院（YashanDB）在 2023 深圳企业创新发展大会主论坛上，与金融、能源和交通等关键行业的 7 家标杆企业（珠海华润银行股份有限公司、深圳市燃气集团股份有限公司、深圳市环境水务集团有限公司、深圳市地铁集团有限公司、深圳市智慧城市科技发展集团有限公司、深圳市长亮科技股份有限公司、金蝶软件（中国）有限公司）签约，签约总金额近亿元，标志着 YashanDB 在重点行业的商用落地迈入全面提速阶段。

- **沃趣科技完成数千万 B+轮战略融资**

2023 年 11 月 23 日，沃趣科技正式官宣，完成了数千万 B+轮战略融资。本轮融资由翌马投资（恒生电子产业基金）领投，金蚂投资跟投，融资金额将用于研发投入及市场拓展。沃趣科技成立于 2012 年，总部位于杭州。针对客户核心业务场景，沃趣科技的数据库专有云产品 - QData 数据库一体机提高了稳定性并实现了近十倍的性能提升。针对多种数据库通用场景，沃趣科技的私有云 RDS 平台，与近 15 款数据库、20+ 云厂商完成产品深度合作，大幅降低企业在多库、多云环境下的运维难度与管理成本。

- **蚂蚁流式图计算引擎 TuGraph Analytics 打破 LDBC 记录**

2023 年 12 月 9 日，国际关联数据基准委员会（简称 LDBC）发布了图数据基准测评“LDBC SNB-BI”最新结果。由蚂蚁集团自研的流式图计算引擎 TuGraph

Analytics 在 30TB 规模的数据集上成功完成了基准测试，数据规模和性能打破了此前的公开纪录，关键指标中的并发吞吐量提升至 2.84 倍，查询能力提升至 1.86 倍。

- **时序数据库 HoraeDB 正式加入 Apache 孵化器**

2023 年 12 月 11 日，蚂蚁集团将时序数据库 CeresDB 的核心源代码以 HoraeDB 的品牌捐赠给 Apache Software Foundation (ASF)。HoraeDB 正式加入 Apache 孵化器。蚂蚁集团时序数据库 CeresDB 于 2022 年 6 月正式宣布开源。HoraeDB 加入 Apache 孵化器后，HoraeDB 将积极践行开放、协作 的 Apache 之道，保持开放治理，持续构建一个公平、多元、包容的社区。

- **OpenTenBase、GreatSQL 等正式捐赠给开放原子基金会**

2023 年 12 月 16 日，2023 开放原子开发者大会在江苏无锡成功举办。大会上，腾讯云将企业级分布式数据库 TDSQL 的社区发行版 OpenTenBase 正式捐赠给开放原子基金会。万里数据库将开源数据库项目 GreatSQL 捐赠给开放原子开源基金会，GreatSQL 荣获“2023 快速成长开源项目”。Apache IoTDB 荣获“2023 生态开源项目”。华为 openGemini 先后获得开放原子基金会授予的“2023 快速成长开源项目”奖项，以及由中国信通院评估颁发的“可信开源项目”和“可信数据库”两项证书。

四、大事记——奖项荣誉

- **2022 年度 CCF 科学技术奖名单公布**

2023 年 2 月 18 日，2022CCF 颁奖典礼在北京隆重举行，中国计算机学会 (CCF) 颁布了 2022 年度“CCF 科技成果奖”。华为与清华大学、中国移动联合申报的“openGauss：企业级开源数据库系统”获得 2022 年度“CCF 科技成果奖”下设的“科技进步特等奖”。人大金仓创始人王珊教授荣获 2022 年“CCF 最高科学技术奖”，OceanBase 首席科学家阳振坤博士获得了 2022 年 CCF 王选奖，其是国内最为权威的衡量企业、高校技术创新水平的科技奖项。

- **腾讯云数据库 8.14 亿登顶 TPC-C 排行榜**

2023 年 3 月 24 日消息，据数据库领域权威测评机构国际事务处理性能委员会 (TPC, Transaction Processing Performance Council) 官网披露，腾讯云数据库 TDSQL 顺利通过 TPC-C 基准测试，性能达到每分钟 8.14 亿笔交易 (tpmC)，打破世界纪录。这也标志着国产数据库 TDSQL 的分布式架构设计和资源调度能力，均达到了业界顶尖水平。

- **OceanBase 与 YashanDB 入选 2023 数字中国建设峰会 “十大硬核科技”**

2023 年 4 月 26 日至 28 日，由国家网信办、国家发改委、科技部、工信部、国务院国资委、福建省人民政府共同主办，福州市人民政府和有关单位承办的第六届数字中国建设峰会在福州举行。单机分布式一体化数据库 OceanBase 4.0（小鱼）、深圳计算科学研究院崖山数据库系统 YashanDB 入选大会“十大硬核科技”。浪潮集团旗下瀚高数据库管理系统入选国资委年度十大技术成果榜单。”达梦启云数据库云服务系统“在 2023 数字中国创新大赛·信创赛道全国总决赛中斩获一等奖。

- **GoldenDB 连续三年入围工信部网安中心 “信息技术应用创新解决方案”**

2023 年 4 月 27 日，由工信部等部门联合主办的数字技术融合创新应用研讨会在福州市举办。本次研讨会重磅发布了 2022 数字技术融合创新应用解决方案征集工作成果，金篆信科旗下 GoldenDB 分布式数据库凭借“面向关键行业重点应用的国产信创分布式数据库解决方案”成功入围“2022 年信息技术应用创新解决方案（典型解决方案）”。值得一提的是，这也是 GoldenDB 分布式数据库连续第三年获此殊荣。

- **金仓数据库、航天天域数据库入选国资委 “中央企业科技创新成果”**

2023 年 5 月 20 日，国务院国资委发布《中央企业科技创新成果推荐目录（2022 年版）》。中国电科成员单位人大金仓旗下金仓云数据库、中国航天科工集团旗下航天天域数据库，以其领先的技术水平和创新成果，在基础软件领域成功入选，再次得到国家的权威认可。

- **GoldenDB 的工商银行案例入选 2023 数博会 “十佳大数据案例”**

2023 年 5 月 27 日，中兴通讯旗下金篆信科 GoldenDB 申报的案例“基于大数据技术的面向金融交易场景的数据库产品在工商银行的应用”成功入围 2023 数博会“十佳大数据案例”。据悉，在该案例中，银行在对公、清算等业务板块引入使用 GoldenDB 分布式数据库，使用范围涵盖了 12 个应用系统、62 个应用租户、2600 多个服务节点、千万级日均交易量。此次获奖，证明了金篆信科 GoldenDB 在大数据领域的实力和成果，也彰显了其在金融交易场景数据库产品方面的优势 and 创新能力。

- **达梦数据荣获 “2022 年度全国版权示范单位” 称号**

2023 年 6 月 16 日，国家版权局公布 2022 年度全国版权示范单位、示范单位（软件正版化）和示范园区（基地）名单，武汉达梦数据库股份有限公司（简称达梦数据）被授予“全国版权示范单位”称号，这是继达梦数据获评“国家企业技术中心”、“博士后科研工作站”后获得的又一国家级殊荣。

- **南大通用 GBASE 荣获 2023 年数交会“创新竞争力产品奖”**

2023 年 7 月 6 日，由商务部、科技部、贸促会和辽宁省政府主办的 2023 中国国际数字和软件服务交易会(简称“数交会”)在大连举办。在大会组织的中国数字和软件服务业评选中,GBASE 南大通用旗下分布式逻辑数仓 GBase 8a MPP Cluster 荣获「创新竞争力产品奖」。此次获奖的分布式逻辑数仓 GBase 8a MPP Cluster 是南大通用自主研发的核心产品，是行业首创的列存数据库，基于 Shared Nothing 架构，面向大数据分析类应用领域，可用作数据仓库系统、BI 系统和决策支持系统的承载数据库。

- **腾讯云 TDSQL 获评“2023 年度可信云金融行业服务最佳实践”**

2023 年 7 月 25 日，中国信息通信研究院在 2023 可信云大会正式发布了 2023 年度可信云最佳实践结果，腾讯云数据库 TDSQL 凭借业内领先的技术和在金融行业数字化转型丰富的实战经验，获评“2023 年度可信云金融行业服务最佳实践”，这代表 TDSQL 推动金融行业数字化转型的一站式解决方案和服务案例受到了行业权威机构的高度认可。

- **2023“金鼎奖”评选结果揭晓**

2023 年 8 月 24 日，由中国人民银行直属企业中国金融电子化集团有限公司主办的 2023 中国国际金融展“金鼎奖”评选结果公布：OceanBase、TiDB、华为 GaussDB、科蓝 SUNDB 荣获 2023 年“金鼎奖”-优秀金融科技解决方案奖。腾讯云 TDSQL、OceanBase、南大通用 Gbase、柏睿数据 RapidsDB、人大金仓荣获 2023 年“金鼎奖”-优秀网信产品基础软件奖。

- **OceanBase 荣获 2023 年软博会银奖**

2023 年 8 月 31 日，中国电子信息行业联合会主办的“第 25 届中国国际软件博览会（简称‘软博会’）”在天津开幕，OceanBase 原生分布式关系数据库荣获“软博会大奖”银奖，是数据库基础软件领域的唯一获奖产品。软博会是我国首个以软件为主题的国家级专业化展会，是我国软件和信息技术服务领域内规模最大、持续时间最长、最具影响力的专业盛会。此次 OceanBase 能入围中国软件行业最高等级评选，不仅是对 OceanBase 完全自主研发分布式数据库产品实力的认可，更是对 OceanBase 坚持长期主义、努力代表中国科技走向全球的肯定。

- **云和恩墨 MogDB、万里数据库 GreatDB 荣获“2023 世界计算大会专题展优秀成果”奖**

2023 年 9 月 15 日，由湖南省人民政府、工业和信息化部联合主办的 2023 世界计算大会在湖南长沙拉开帷幕。云和恩墨“MogDB 数据库金融行业国产化实践方案”和万里安全数据库 GreatDB“事务型关系数据库解决方案”获评“2023 世界计算大

会专题展优秀成果”奖。达梦数据入选 “2023 中国先进计算企业百强榜” TOP50，也是国产数据库领域唯一入围企业。

- **OceanBase 再获 OSCAR 两项大奖，坚定开源开放**

2023 年 9 月 21 日，在由中国信息通信研究院主办的“2023 OSCAR 开源产业大会”上，蚂蚁集团旗下的自研原生分布式数据库 OceanBase 荣获“2023 OSCAR 尖峰开源项目”、“2023 OSCAR 尖峰开源企业（开源运营与生态建设）”两个奖项。同时，完成可信开源社区评估，获得“可信开源社区”评估证书。

- **人大金仓与达梦获评“2023 年度软件和信息技术服务名牌企业”**

2023 年 10 月 13 日，中国电子信息行业联合会在 2023 世界数字经济大会暨第十三届智慧城市与智能经济博览会上宣布，北京人大金仓信息技术股份有限公司与武汉达梦数据库股份有限公司获评“2023 年度软件和信息技术服务名牌企业”。亚信科技（AntDB）数据库再获殊荣，获评“2023 年信创物联网优秀服务商”。此次评选由工信部信发司指导，旨在贯彻落实国家《“十四五”软件和信息技术服务发展规划》，培育行业品牌，提升企业市场竞争力，进一步推动我国软件产业做强做优做大。

- **刘钰博士获颁 2022 年度“CCF-腾讯犀牛鸟基金”优秀专利奖**

2023 年 10 月 26 日，第二十届中国计算机大会 CNCC 2023 年在沈阳隆重举办。北京交通大学计算机与信息技术学院刘钰博士获颁 2022 年度“CCF-腾讯犀牛鸟基金”优秀专利奖。刘钰博士的科研基金题为“面向图视角的多模数据库高效查询与计算技术研究”，该项目是 2022 年度 CCF-腾讯犀牛鸟基金项目中唯一获批的数据库领域的研究课题。

- **gStore、CeresDB、GreptimeDB 、 RisingWave 入选“2022-2023 年度世界开源贡献榜”**

2023 年 11 月 7 日，Bench Council（国际测试委员会）公布了“2022-2023 年度世界开源贡献榜”，该榜单号称“只以贡献分高下”。中国的国产图数据库 gStore、时序数据库 CeresDB、GreptimeDB 与 流数据库 RisingWave 入选“2022-2023 年开源系统杰出成果”；北京大学邹磊作为 gStore 主要学术贡献者入选“开源系统杰出人才”。

2023 年 BenchCouncil（国际测试委员会）邀请多位独立科学家，从 2022 至 2023 年度的数万项开源相关成果中，遴选出了 102 项代表性成果，在确定主要贡献者的基础上产生了开源领域年度人才榜、机构榜、国家榜。共 195 人进入榜单。美国在国家榜上领先（43.5 分），中国紧随其后排名第二（37 分）。

- **华为云 GaussDB 荣获“2023 世界互联网大会领先科技奖”**

2023 年 11 月 8 日，2023 年世界互联网大会乌镇峰会隆重举行。在大会的“世界互联网大会领先科技奖颁奖典礼”上，华为云 GaussDB 数据库凭借领先的技术优势成功获得“2023 世界互联网大会领先科技奖”。这是行业对 GaussDB 数据库技术实力的高度认可。在技术能力的背后，是华为云 GaussDB 全球 700 多项核心专利、80 多篇顶会论文、2000 多名内核研发人才的强有力支撑。如今，华为云 GaussDB 已经在华为终端云上线了 6000 多个分布式数据库节点，数据管理规模超 6PB，服务了 4 亿多用户。

- **金篆信科 GoldenDB 荣获高交会“优秀产品奖”**

2023 年 11 月 15 日-19 日，由商务部、科学技术部、工业和信息化部、国家发展改革委、农业农村部、国家知识产权局、中国科学院、中国工程院和深圳市人民政府共同举办的第二十五届中国国际高新技术成果交易会（简称“高交会”）顺利举办。金篆信科 GoldenDB 亮相本届大会，并荣获“优秀产品奖”！此次获得优秀产品奖，代表着高交会对金篆信科 GoldenDB 技术实力、产品能力、服务水平、市场影响力的认可。在数据库最为关键的内核引擎，GoldenDB 目前已经实现了全自研；支持多模 SQL 引擎，支持更多类型数据的存储和使用、支持智能负载分析和调整，优化资源利用。

- **工信部“2023 年中国赛宝信创优秀解决方案”奖项公布**

2023 年 11 月 23 日，由工业和信息化部电子第五研究所和中国通信企业协会联合主办的“第二届中国赛宝信息技术应用创新优秀解决方案征集活动”结果正式公布。星环科技基于湖仓一体架构的新一代医院数据中心解决方案荣获“应用创新示范方向一等奖”。云和恩墨基于 MogDB 的“新一代医疗 HIS 系统全栈 XC 解决方案”获颁“应用创新示范方向二等奖”。创邻科技“Galaxybase 智能风控解决方案”荣获“全栈解决方案全国三等奖”！万里数据库“基于 GreatDB 的全栈信创数据库解决方案”荣获“全栈解决方案优秀奖”。

- **国家工信安全中心“2023 年数字化转型自主创新解决方案优选案例”公布**

2023 年 11 月 28 日，国家工业信息安全发展研究中心发布“数智赋能，创新领航”——2023 年数字化转型创新解决方案优选计划评审结果。上海云熹“基于 KaiwuDB 的分布式储能行业解决方案”、深圳计算科学研究院“崖山数据库系统融合数据管理解决方案”、优炫软件“基于优炫数据库的银行密钥管理平台的构建与实现”、达梦数据““互联网+北京交通”综合监管平台服务”、东方通“分布式数据库缓存解决方案”、东方国信“湖仓一体解决方案”、力控元通“化工行业实时数据库国产化方案”等获评“数字化转型创新解决方案优选案例”。

- **2023 年大数据“星河案例”之数据库标杆案例及优秀案例发布**

2023 年 12 月 20 日，在 2023 数据资产管理大会上，中国信息通信研究院、中国通信标准化协会大数据技术标准推进委员会(CCSA TC601)共同发布了“2023 大数据“星河”案例”。14 个数据库标杆案例（KaiwuDB、AntDB、Kingwow、磐维等），26 个数据库优秀案例（达梦、HotDB、Apache Doris、磐维等）入选。

五、大事记——版本发布

（一）国外数据库重要产品发布：

- **时序数据库 InfluxDB 3.0 发布**

2023 年 4 月 26 日，InfluxData 发布了 InfluxDB 3.0。InfluxDB 3.0 在供应商的云产品——InfluxDB Cloud Serverless 和 InfluxDB Cloud Dedicated 中普遍可用，并将在今年晚些时候在供应商的自我管理产品中提供。新版本不仅旨在改进供应商平台的先前版本，而且还作为未来添加到 InfluxDB 的基础。

- **PostgreSQL 16 版本正式发布**

2023 年 9 月 14 日，PostgreSQL 全球开发集团宣布发布 PostgreSQL 16，这是世界上最先进的开源数据库的最新版本。PostgreSQL 16 包含许多新功能和增强功能，包括：允许使用 OUTER 和 FULL 连接并行执行查询、允许从备用服务器进行逻辑复制、允许逻辑复制订阅者并行应用大事务等功能。此版本为开发人员和管理员提供了许多功能，包括更多 SQL/JSON 语法、针对工作负载的新监控统计数据，在管理大型集群的策略的访问控制规则定义方面，有更大的灵活性。

（二）国产数据库重要产品发布：

- **浪潮 KaiwuDB 1.0 版本发布**

2023 年 1 月 10 日，浪潮集团控股的 KaiwuDB 正式推出 1.0 版本。KaiwuDB 1.0 由 KaiwuDB 团队核心自主研发，面向工业物联网、数字能源，车联网、智慧产业等行业场景；拥有“就地计算”等专利技术，具备每秒百万级写入、千万条记录聚合查询毫秒级响应，满足海量、高并发的时序数据写入及快速查询和复杂查询的需求；同时具备集群部署、高可用、低成本、易运维等特性。

- **蚂蚁集团发布 CeresDB 1.0，Rust 高性能云原生时序数据库**

2023 年 3 月 1 日，蚂蚁集团云原生时序数据库 CeresDB 1.0 正式发布，达到生产可用标准。CeresDB 是一款高性能、分布式的云原生时序数据库，采用 Rust 编写，

与经典时序数据库相比，CeresDB 的目标是能够同时处理时序型和分析型两种模式的数据，并提供高效的读写。

- **openGauss 5.0.0 里程碑版本发布**

2023 年 3 月 31 日，openGauss 5.0.0 里程碑版本发布！openGauss 5.0.0 是 openGauss 发布的第三个 LTS 版本，版本生命周期为 3 年。openGauss 5.0.0 版本与之前的版本功能特性保持兼容，在内核能力、工具链、兼容性方面全面增强。

- **瀚高发布瀚高数据库企业版 V9 等新品**

2023 年 4 月 7 日，由瀚高数据库主办的“2023 瀚高生态大会”在青岛成功举办。瀚高股份副总裁吕新杰博士现场发布了 IvorySQL 开源根社区以及瀚高数据库企业版 V9、隐私计算数据库等新品，进一步加强了瀚高数据库全栈产品体系，拓宽了数据全生命周期管理范围，为数据库产业再添创新利器。

- **Zilliz Cloud GA 版本正式发布**

2023 年 4 月 7 日，Zilliz Cloud GA 版本正式发布！本次 Zilliz Cloud 大版本更新提升了 Zilliz Cloud 向量数据库的可用性、安全性和性能，并推出了一系列新功能。

- **庚顿正式发布军工版实时数据库庚金 3.0**

2023 年 4 月 21 日消息，庚顿正式发布军工版实时数据库庚金 3.0，鼎力支撑中国国防数字化。庚金 3.0 充分研究实时处理技术和数据管理技术，融合领先的实时操作系统、自动化系统、安全防护系统、关系型数据库、NoSQL 数据库、计算引擎等相关技术，具有高可靠、高可用、高安全、跨平台、可伸缩的特性，结合庚顿自主研发的数据处理产品可实现对海量时序数据的采集、存储、计算、传输、展现、分析等各个环节的高效管理，特别适用于军工、航天等国家重要领域。

- **万里数据库发布新一代数据库一体机**

2023 年 5 月 8 日，华为中国合作伙伴大会 2023 在深圳举行。现场，创意信息联合华鲲振宇、拓林思软件、万里数据库发布新一代数据库一体机，吸引众多行业嘉宾关注。

- **爱可生发布基于 OceanBase 的商业发行版 ActionDB**

2023 年 5 月 19 日，上海爱可生信息技术股份有限公司召开数据库新品发布会，正式发布面向国产化时代的企业级数据库 ActionDB。ActionDB 是基于 OceanBase 开源内核的商业发行版，是一款高性能、高可用、高扩展的单机分布式一体化数据库产品。ActionDB 继承了 OceanBase 稳定可靠、高性能的优点，并充分发挥爱可生多年

在开源数据库领域的专业经验和技術积累，提供了企业级的安全特性、易用的数据库工具和更便捷的服务。

- **openGauss Developer Day 2023，多伙伴联合发布 openGauss 数据库一体机**

2023 年 5 月 25 日 - 26 日，openGauss Developer Day 2023（openGauss 开发者大会 2023）在北京举办。会上，openGauss 推出 DataPod+DataKit 组合和第三代智能优化器 ABO，打造全新的数据底座；海量数据、云和恩墨、南大通用、沃趣科技正式发布首批基于 openGauss 发行版的数据库一体机产品；云和恩墨创始人盖国强、华为鲲鹏计算业务总裁李义及多家 DBV 伙伴代表联合发布基于 openGauss 5.0 的数据库商业发行版。openGauss 社区正式发布 openGauss 伙伴专业保障服务，首批 8 家优秀伙伴通过认证成为社区服务伙伴；云和恩墨 MogDB 数据库荣获“鲲鹏最佳基础软件奖”这是 MogDB 连续第二年获得该殊荣；此外，云和恩墨数据库内核研发工程师王修强获颁“2022 年度鲲鹏优秀开发者”称号；随着 openGauss5.0 内核发布，GBASE 发布新品 GBase 8c V5 和 GBase 8s V8.8.5 版本。

- **星环科技发布多个数据库新版本和向量数据库 Transwarp Hippo**

2023 年 5 月 26 日，由星环科技、上海数据交易所、上海大数据联盟、财联社联合主办的向星力·未来数据技术峰会（FDTC）在上海举办。会上，星环科技众多分布式数据库产品发布新版本（能够替代国外产品的分布式分析型数据库 ArgoDB 5.0、分布式交易型 KunDB 3.2；构建海量数据互联智慧“星”图的分布式图数据 StellarDB 5.0；面向多元场景的高性能时序数据库 TimeLyre 9.1 等），并推出了分布式向量数据库 Transwarp Hippo，拓展大语言模型时间和空间维度，高效地解决向量相似度检索以及高密度向量聚类等问题。

- **Milvus 社区发布 Milvus Lite 轻量级版本**

2023 年 6 月 8 日，Milvus 社区发布 Milvus Lite，即 Milvus 的轻量级版本，方便有相关需求的用户进行体验。Milvus Lite 为没有专业运维团队支撑、安装部署环境受限的群体提供了新的可能。

- **腾讯云发布向量数据库 Tencent Cloud VectorDB**

2023 年 7 月 4 日，腾讯云正式发布国内首个 AI 原生（AI Native）的向量数据库 Tencent Cloud VectorDB。其最高支持业界领先的 10 亿级向量检索规模，并将延迟控制在毫秒级。相比传统单机插件式数据库检索规模提升 10 倍，同时具备百万级每秒查询（QPS）的峰值能力。

- **云和恩墨 MogDB 5.0.0 和 Uqbar 2.0.0 开放下载**

2023 年 7 月 14 日，云和恩墨旗下企业级关系型数据库 MogDB 5.0.0 和超融合时序数据库 Uqbar 2.0.0 正式在官网上线并开放下载。早前，在 openGauss Developer Day 2023 云和恩墨专题论坛上二者已正式发布。企业级关系型数据库 MogDB 5.0.0 作为 LTS 版本，基于 3.0/3.1 版本进一步增强，并引入了 openGauss 5.0.0 的全部特性；同时增加多项自研特性。超融合时序数据库 Uqbar 2.0.0 是首个 LTS 版本，其继承了 1.1.0 的功能，具备时序生命周期策略、时序表管理、数据压缩、过期删除、持续聚合管理、高可用、备份与恢复、运维工具等基本功能。

- **沃趣科技联合超聚变发布“QData T7 数据库一体机解决方案”**

2023 年 8 月 18 日，2023 超聚变合作伙伴大会在北京举行，沃趣科技联合超聚变发布了“QData T7 数据库一体机解决方案”，通过端到端全栈软硬协同，在 OLAP 场景应用上，性能相较上一代产品提升 2.5 倍，达到业界全球领先水平。沃趣科技与超聚变联合打造的 QData T7 数据库一体机，基于超聚变最新一代 FusionServer 2288H V7 服务器，将服务器、网络、存储进行高度集成，实现业务性能、平台可靠性、业务连续性及可用性的提升。

- **Apache Doris 2.0.0 版本正式发布，盲测性能 10 倍提升**

2023 年 8 月 11 日，Apache Doris 2.0.0 版本正式发布。在 2.0.0 版本中，Apache Doris 在标准 Benchmark 数据集上盲测查询性能得到超过 10 倍的提升、在日志分析和数据湖联邦分析场景能力得到全面加强、数据更新效率和写入效率都更加高效稳定、支持了更加完善的多租户和资源隔离机制、在资源弹性与存算分离方向踏上了新的台阶、增加了一系列面向企业用户的易用性特性。

- **南大通用重磅发布国产向量数据库 GBase Cloud Vector DB**

2023 年 9 月 1 日，由天津市工业和信息化局指导，天津南大通用数据技术股份有限公司承办的第二十五届中国国际软件博览会·中国数据库产业峰会在天津顺利召开。会上，GBASE 南大通用重磅发布国产向量数据库 GBase Cloud Vector DB。GBase Cloud Vector DB 在国内领先的 GBase 8a 集群基础上实现，可以被广泛应用于各类 AI 驱动的应用场景，为大模型关键技术和极致数字世界而来，代表了当今数据库技术的最新发展趋势和成果，向量数据库的推出，进一步完善了 GBASE 全栈的数据库产品矩阵。

- **腾讯云正式发布 TDSQL 融合版**

2023 年 9 月 7 日，在 2023 腾讯全球数字生态大会腾讯云数据库 TDSQL 技术与实践专场，腾讯云正式发布 TDSQL 融合版，该版本整合了之前 TDSQL 系列产品内核的优势，并在内核架构、Oracle 兼容能力、性能、隔离、迁移工具等多个关键能力

上进行了大幅增强优化，为企业打造了更加极致的 HTAP 国产精品数据库。值得一提的是，新升级的 TDSQL-C 实现了全球首个可释放存储架构的 Serverless 服务。

- **2023 亚信科技 AntDB 数据库 8.0 产品发布，“超融合+流式实时数仓”**

2023 年 9 月 20 日，亚信科技 AntDB 数据库 8.0 产品发布会成功举办。AntDB 数据库 8.0 产品实现了两大特性的重磅升级：“超融合架构”从实验室走向生产，流式计算升级为“超融合流式实时数仓”。

- **华为云正式发布 GaussDB 向量数据库，并联合邮储银行发布企业级数据库 Ubase**

2023 年 9 月 20 日至 22 日，华为全联接大会 2023（HC2023）在上海盛大召开。华为云正式发布 GaussDB 向量数据库。GaussDB 向量数据库基于 GaussDB 开发，具备一站式部署、全栈自主可控的优势，并且在 ANN-Benchmarks 中排名第一，技术实力深厚。在大模型技术、产品和应用层出不穷的当下，GaussDB 向量数据库将为大模型行业深度赋能，加速盘古大模型行业落地。9 月 22 日，openGauss 社区联合邮政储蓄银行正式发布首个基于 openGauss 的金融行业企业自用版数据库 Ubase。

- **SelectDB 全新产品形态全面发布，定义现代化实时数据仓库**

2023 年 9 月 25 日，2023 飞轮科技产品发布会在线上召开。飞轮科技 CEO 马如悦全面解析了现代化数据仓库的演进趋势，宣布立足于多云之上的 SelectDB Cloud 云服务全面开放，增加了全新的私有仓库（BYOC）产品模式，同时发布了更加自主可控的 SelectDB Enterprise 企业版。

- **蚂蚁开源图数据库 TuGraph-DB 升级到 v4.0，全新支持 GQL 国际标准查询语言**

2023 年 10 月 9 日，蚂蚁开源图数据库 TuGraph-DB 升级到 v4.0。TuGraph-DB v4.0 主要支持了 ISO GQL 国际标准查询语言、企业级高可用能力以及图学习引擎。TuGraph-DB v4.0 新版本遵循 GQL 的最新标准，在查询引擎进行了支持，为用户提供了丰富多样的查询语言选择，也对图数据库领域查询语言的标准化起到了推动作用。TuGraph-DB 开源了企业级高可用能力，能够实现多活热备。TuGraph-DB v4.0 深度集成了图学习引擎，兼容 DGL、PyG 等常见图学习框架。

- **“海量数据 1024 开发者日”发布 Vastbase G100 V3.0**

2023 年 10 月 24 日，2023 年海量数据 1024 开发者日暨产品发布会圆满落幕。会上，海量数据相继发布海量数据库 Vastbase G100 的最新 3.0 版本、Vastcube G3000 数据库一体机、Datalink-PLAT 企业级大数据应用平台，正式官宣海量数据自主产品体系，着力释放数据要素价值，推动数据全产业链生态、良性发展。Vastbase G100

V3.0 带来了重要的创新和突破：一是实现了资源池化、一写多读的创新架构，二是推出了超融合兼容引擎，三是面向时序场景新增了时序引擎。

- **东方国信重磅发布 CirroData for MySQL V2.0 新一代实时数数仓**

2023 年 10 月 27 日，东方国信重磅发布了 CirroData for MySQL V2.0 新一代实时数仓。CirroData for MySQL 拥有卓越的性价比、超高查询速度、融合统一、简单易用、方便运维等特点，以应对企业在数字化转型中面临的海量数据、实时分析、敏捷开发等挑战。

- **YashanDB V23.1 三大新品重磅发布，核心场景再突破**

2023 年 11 月 8 日，YashanDB 2023 年度产品发布会在线上成功召开，YashanDB V23.1 版本发布。面向企业级核心应用、大数据分析、空间数据管理场景，YashanDB 首次发布 YashanDB for Cluster 共享集群、YashanDB for Data Warehouses 分布式实时数仓以及 YashanDB for GIS 空间数据库三大产品形态。YashanDB 个人版已正式向所有用户和开发者全面开放下载。

- **华为鲲鹏联合 19 家合作伙伴共同发布基于鲲鹏底座的一体机产品**

2023 年 11 月 13 日，2023 鲲鹏一体机技术应用大会在北京召开。华为鲲鹏联合 19 家合作伙伴共同发布基于鲲鹏底座的一体机产品。云和恩墨发布基于鲲鹏 CPU 架构服务器和 openEuler 操作系统的数据库一体机产品 zData X。该产品能够支持包括 openGauss、MogDB、达梦、人大金仓等在内的多种国产数据库，满足企业全栈国产建设要求。海量数据联发布基于国产硬件平台和自主可控数据库 Vastbase 打造的一体机产品 Vastcube G3000 数据库一体机。该产品具备快速全栈国产化替代能力；同时具备高性能、高可靠；云化管理，快速部署，开箱即用等特点，能够满足客户多场景应用需求。柏睿数据与华鲲振宇、鲲鹏联创联合发布“智算一体机”。柏睿智算一体机融合计算、存储、网络、数据库、监控于一体，以全内存分布式计算引擎 RapidDB 和向量引擎 Rapids VectorDB 为核心，支持全场景的海量多模态数据高效管理与实时分析。

- **2023 OceanBase 年度发布会发布一体化数据库 OceanBase 4.2.1 LTS**

2023 年 11 月 16 日，2023 OceanBase 年度发布会在京顺利召开。在本次发布会上，OceanBase 对外正式宣布：将持续践行“一体化”新战略，为关键业务负载打造一体化数据库。同时，会上发布一体化数据库的首个长期支持版本 OceanBase 4.2.1 LTS，标志着 OceanBase 一体化数据库进入可规模化上线使用的长期支持阶段。OceanBase 4.2.1 LTS 是面向 OLTP 核心场景的里程碑版本，具备 OLTP 的完整功能，相比 3.2 版本有更强的 OLTP、OLAP 性能，相比传统容灾提供更具性价比的仲裁无损容灾方案，可通过 2 个副本实现 RPO=0。

- **瀚高数据库发布 IvorySQL 3.0，基于 PG 16.0 最新内核**

2023 年 11 月 16 日，瀚高数据库发布 IvorySQL 3.0。此版本不仅继承了 PostgreSQL 16.0 的最新内核和功能，还扩展了更多企业级特性。相比于 PostgreSQL 社区版，IvorySQL 3.0 在兼容性和易用性方面实现了显著提升，同时为适应容器化和云端环境提供了更为全面的支持。3.0 具备更为完善的数据库特性和创新的功能，提供高度的 SQL 和 PL/SQL 兼容性。

- **虚谷数据库 V12.0 发布，聚焦信息创新与数字转型**

2023 年 11 月 25 日，由信创工委联合江苏省工信厅、常州市人民政府共同举办的 2023 信息技术应用创新论坛在常州开幕。虚谷伟业在大会发布了虚谷最新产品——虚谷数据库 V12.0，虚谷数据库新一代产品聚焦信息创新与数字转型，提出数字化与信息创新双轮驱动。虚谷数据库兼容了 Oracle 和 MySQL 开发者生态，包括 SQL 语法、数据类型、关键字以及 Oracle 的存储过程语法。

- **TiDB 7.5 LTS 发版，提升规模化场景下关键应用的稳定性和成本的灵活性**

2023 年 12 月 1 日，TiDB 7.5 LTS 发版。作为 TiDB 7 系列的第二个长期支持版本 (LTS)，TiDB 7.5 着眼于提升规模化场景下关键应用的稳定性。新版本中，TiDB 在可扩展性与性能、稳定性与高可用、SQL 以及可观测性等方面获得了持续提升。

- **2023 IoTDB Summit 用户发布 IoTDB 企业版 V1.3**

2023 年 12 月 3 日，2023 IoTDB 用户大会在北京成功举行。在峰会主论坛中，天谋科技的 CTO、Apache IoTDB PMC Member 乔嘉林正式发布了 IoTDB 企业版 V1.3。新发布的 IoTDB 企业版涵盖多个内核技术亮点，包括三个工具 (IoTDB - OpsKit 部署工具、系统监控工具、可视化控制台工具)、两个引擎 (流处理引擎、智能分析引擎)、与一个项目 (Apache TsFile，时序数据标准文件格式，是继 IoTDB 后，时序数据领域第二个 Apache 顶级项目)，实现了“单平台采存算管用”的横向一站式解决方案，与“跨平台端边云协同”的纵向一站式解决方案。

- **百度发布自研云原生数据库 GaiaDB 4.0，增强并行查询能力**

2023 年 12 月 20 日，2023 百度云智大会·智算大会在北京举办。会上，百度发布自研云原生数据库 GaiaDB 4.0，该数据库增强了并行查询能力，突破单机计算瓶颈，实现跨机多核并行查询，在混合负载和实时分析业务场景中性能提升超过 10 倍。

10、数据库国产化进程及标准化政策

一、信创政策持续催化，加速国产数据库发展

● 信创产业发展历程：

1. 萌芽期 (1999-2005): 以“缺芯少魂”为警示，Xteam、蓝点等国产厂商陆续成立，为信创产业奠定基础。
2. 起步期 (2006-2013): “核高基”政策出台，明确发展方向，信创产业起步。
3. 试点实践期 (2014-2017): 中央网信办成立，各地部署信创试点，产业生态初步形成。
4. 规模化推广期 (2018-至今): 中兴/华为事件加速信创产业发展，应用范围不断扩大，2020 年被视为规模化应用元年。

● **国产厂商破局崛起，开启自主创新之路**”。数据库技术历经 80 余载的演变，从最初的单一形态发展到如今百花齐放的局面。关系型数据库依然占据市场主导地位，但新型数据库如雨后春笋般涌现，为数据库技术发展注入了新的活力。纵观全球数据库市场，国外厂商凭借先发优势长期占据主导地位。然而，近年来国产数据库厂商奋起直追，在技术创新、政策支持和市场需求的共同推动下，实现了弯道超车，打破了国外厂商的垄断局面。技术创新是国产数据库破局的关键。国产数据库厂商积极拥抱新技术，在 NoSQL、NewSQL、分布式等领域取得了突破性进展，并在性能、可靠性、易用性等方面不断提升，与国外厂商的差距逐渐缩小。

● **数据库厂商迅速抢占市场份额**。全球数据库市场竞争激烈，各厂商各显神通，纷纷抢占市场份额。在本地部署的关系型数据库领域，Oracle 等国外厂商凭借先发优势仍占据一定优势。然而，近年来，中国数据库厂商在政策支持和自身技术实力的双重推动下，实现了快速发展，市场份额连年增长，2021 年已达 64%。值得注意的是，中国数据库厂商的营收中，相当一部分来自托管国外开源数据库，自主创新能力仍需进一步提升。未来，国产数据库厂商需加大技术创新力度，打造自主知识产权的产品，才能在竞争中站稳脚跟，实现国产数据库的全面替代。

● **数据库生态与服务：补短板，强优势**。国产数据库起步较晚，在传统领域追赶国外厂商仍需时间。然而，新技术、新场景为国产数据库带来了弯道超车的机遇。信创数据库厂商应抓住机会，在产品方面努力追赶，在生态与服务方面补齐短板，建立综合优势。生态与服务建设是数据库发展的重要基石。目前，国产数据库生态与服务尚不完善，相关机制和伙伴需要长期建设。只有打造良好的生态环境，才能吸引更多用户和开发者，推动国产数据库的广泛应用。

二、重点测评报告出炉，自主可控取得显著进展

2023 年数据库行业重点测评报告是指由国内外权威机构和企业发布的，对数据库产品进行测评的报告。这些报告主要考察数据库产品的功能性、兼容性、可靠性、性能、安全性等方面。2023 年数据库行业重点测评报告的发布，为用户选择数据库产品提供了参考依据，也促进了数据库产品的技术创新和应用推广。总体来看，2023 年数据库行业重点测评结果表明，自主可控数据库在技术水平、应用场景、安全性等方面取得了显著进展，正在逐步满足国家数据库安全可控的要求。

表 7 2023 年发布数据库重点评测

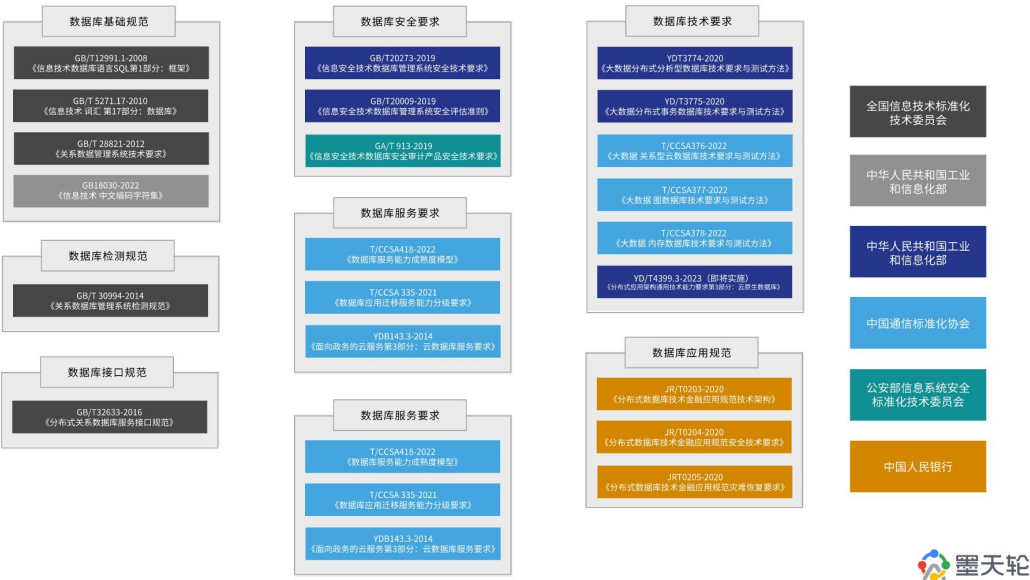
发布日期	评测机构	评测项目	主要内容
2023.6.25	中国信息通信研究院	2023 年上半年“可信数据库”评测	共计 28 家企业的 33 款产品通过本次评审，通过基础能力评测的包括 6 款搜索型数据库、5 款分布式事务型数据库、3 款分布式分析型数据库、3 款全密态数据库、2 款时序数据库、2 款时空数据库、2 款图数据库、1 款关系型云数据库和 1 款数据库一体机；通过性能评测的包括 3 款分布式分析型数据库、2 款集中式事务型数据库、1 款分布式分析型数据库（大规模）、1 款分布式事务型数据库和 1 款时序数据库；通过稳定性评测的包括 1 款分布式分析型数据库和 1 款分布式事务型数据库。此外，通过分布式数据库运维管理能力成熟度模型评估的共 1 家企业。
2023.11.22	中国信息通信研究院	2023 年下半年“可信数据库”评测	2023 年 11 月 21 日和 11 月 22 日，中国信息通信研究院（以下简称“中国信通院”）2023 年下半年“可信数据库”测试的专家评审会圆满结束，共计 25 家企业的 31 项测试通过了本次评审。本次通过基础能力专项测试的包括 9 款向量数据库、3 款时序数据库、2 款分布式分析型数据库、2 款分布式事务型数据库、1 款分布式分析型数据库（大规模）、1 款事务分析混合型（HTAP）数据库、1 款全密态数据库、1 款时空数据库、1 款 Serverless 事务型数据库、1 款搜索型数据库、1 款数据库一体机、1 款数据库管理平台 and 2 款数据库迁移工具；本次通过性能专项测试的包括 1 款分布式分析型数据库、1 款集中式事务型数据库和 1 款时序数据库；

			本次通过稳定性能力专项测试的包括 1 款时序数据库；本次通过数据库服务能力成熟度模型评估的包括 1 家数据库服务提供方。
2023.12.26	中国信息安全测评中心 国家保密科技测评中心	安全可靠测评	安全可靠测评主要面向计算机终端和服务器搭载的中央处理器(CPU)、操作系统以及数据库等基础软硬件产品，通过对产品及其研发单位的核心技术、安全保障、持续发展等方面开展评估，评定产品的安全性和可持续性，实现对产品研发设计、生产制造、供应保障、售后维护等全生命周期安全可靠性综合度量和客观评价。在本次评测中，11 款集中式数据库通过第一批“安全可靠测评”。

三、规范化、标准化行业标准，助力国产数据库健康发展

行业标准准则可以规范行业发展，促进国产数据库产品的健康发展。通过制定标准，明确了数据库产品的功能、性能、安全性等方面的要求，为用户选择数据库产品提供参考依据。在提高产品竞争力，促进技术创新及保障数据安全方面同样具有广泛意义。

截至 2024 年 2 月份，国内现行数据库标准分类及归口单位如下图所示：



图：国内数据库标准分类（关系型）及归口组织

表 8 2023 年发布的数据库行业标准准则

发布时间	发布机构	标准名称	主要内容
2023.1.30	中国信息通信研究院	《数据库应用迁移服务能力分级要求》	共计 28 家企业的 33 款产品通过本次评审，通过基础能力评测的包括 6 款搜索型数据库、5 款分布式事务型数据库、3 款分布式分析型数据库、3 款全密态数据库、2 款时序数据库、2 款时空数据库、2 款图数据库、1 款关系型云数据库和 1 款数据库一体机；通过性能评测的包括 3 款分布式分析型数据库、2 款集中式事务型数据库、1 款分布式分析型数据库（大规模）、1 款分布式事务型数据库和 1 款时序数据库；通过稳定性评测的包括 1 款分布式分析型数据库和 1 款分布式事务型数据库。此外，通过分布式数据库运维管理能力成熟度模型评估的共 1 家企业。
2023.8.1	中国电子技术标准化研究院	GB 18030-2022《信息技术中文编码字符集》	该标准适用范围是具备中文信息处理及交换功能的软硬件产品，设立三档实现级别，共收录汉字 87887 个，比上一版增收了 1.7 万余个生僻汉字。数据库软件对强制性国家标准 GB 18030 的支持程度，直接关系到信息系统的文字处理能力，与生僻字人群的切身利益息息相关。据检测与认证机构统计，我国数据库厂商已有 23 家首次通过测评认证。
2023.8.7	中国信息通信研究院	《数据库迁移工具能力要求》	首批中国信通院“可信数据库”数据库迁移工具专项测试正式启动。中国信通院云计算与大数据研究所依托中国通信标准化协会大数据技术标准推进委员（CCSA TC601），联合 30 余家国内外知名研究院和企业共同完成了《数据库迁移工具能力要求》。《数据库迁移工具能力要求》依托学术界与产业界的深度融合，该标准包含基础功能、数据库对象迁移能力、数据迁移能力、数据校验能力以及迁移评估能力五大能力域，共 44 个测试项，包括 16 个必选项和 28 个可选项。本标准作为业内首个数据库迁移工具的技术标准，可为数据库迁移工具的研发、测试以及选型提供参考。
2023.9.11	中国工业互联网研究院	《国家工业互联网大数据中心体系数据库表命名和设计规范》《国家工业互联网大数据中心体系产业链编码规范》《国家工业互联网大数据中心体系企业资质库数	该批标准明确了业务系统的数据库表命名设计规范，统一了产业链及其上下游节点的编码规范，规定了企业资质、司法风险、产品服务、投融资、知识产权等常见专题库的数据字段，将有助于降低数据对接成本、提升数据质量、提高开发利用效率，为构建一体化的数据体系奠定基础。

		据标准与共享接口规范》	
2023.12.26	财政部、工信部	《数据库政府采购需求标准（2023年版）》	一、采购人采购数据库应当按照《需求标准》实施相关采购活动。二、对于既包含数据库、服务器等硬件产品也包含集成服务的采购项目，采购人应当合理划分采购包，尽可能将数据库、服务器等硬件产品与集成服务分包采购。采购的数据库、服务器等硬件产品总额达到分散采购限额标准的，应当单独分包采购。三、采购人应当加强采购需求管理，按照《政府采购需求管理办法》（财库〔2021〕22号）要求，结合具体应用场景，根据《需求标准》确定采购需求，明确所需数据库的功能、质量等指标要求，并据此编制采购文件采购人应当将《需求标准》中加“”的指标纳入采购需求，并作为采购文件中的实质性要求。其中，乡镇以上党政机关，以及乡镇以上党委和政府直属事业单位及部门所属为机关提供支持保障的事业单位在采购数据库时，应当将数据库符合安全可靠测评要求纳入采购需求，其他单位可不在采购需求中提出此项要求。对于未加“”的指标，采购人可以根据实际需要自行确定是否纳入采购需求。采购人在采购需求中，可以对《需求标准》中的指标提出更高要求，也可以根据实际需要增加《需求标准》以外的指标，但不得超出实际需要。四、供应商在投标、响应环节出具关于所提供数据库满足采购文件要求承诺函的，即视为相关产品符合要求。采购人在供应商投标、响应环节不得对数据库进行检测、认证，也不得要求供应商提供检测报告、认证报告。五、采购人应加强履约验收管理，按照采购合同约定对供应商提供的数据库进行验收，必要时委托依法取得检测、认证资质的机构进行检测、认证。对于供应商未按合同约定提供数据库的，采购人应当依法追究其违约责任。

11、数据库大会展现产业未来图景

近年来，数据库行业大会在国内外得到了广泛关注，参与人数和影响力不断扩大。这表明数据库行业大会在数据库行业中发挥着越来越重要的作用。

- **促进技术交流与传播：**数据库行业大会为数据库技术人员提供了一个良好的交流学习平台，促进了数据库技术的交流与传播。大会上，来自全国各地的数据库专家、学者、企业家将分享最新的研究成果和实践经验，为数据库技术人员提供了学习新知识、拓展新视野的机会。

- **推动行业合作与创新**：数据库行业大会为数据库厂商提供了一个良好的合作交流平台，促进了行业合作与创新。大会上，数据库厂商可以展示自身产品和技术，与其他企业进行交流合作，共同推动数据库产业的发展。

- **引领行业发展**：数据库行业大会为数据库行业的发展指明了方向，推动了行业创新发展。大会上，数据库专家、学者发表了很多主题演讲，分享对数据库行业发展的前瞻性思考，为数据库行业的发展提供新的思路 and 方向。

一、数据库行业大会

- **第十二届数据技术嘉年华(DTC 2023)**

大会时间：2023.4.8-4.9

主办方：中国数据库联盟(ACDU)、墨天轮社区

主要内容：4月8日下午，为期两天的第十二届数据技术嘉年华（DTC 2023）在北京新云南皇冠假日酒店圆满落下帷幕。大会得到了工业和信息化部电子五所的支持和指导，围绕“开源·融合·数字化——引领数据技术发展，释放数据要素价值”这一主题，通过一场主论坛和十二场专题论坛，汇聚“产学研”各界数据技术领军人物、学术精英、技术专家、行业用户，从多角度、多维度带来68场主题演讲。群贤毕至、俊采星驰，从技术发展到方案构建、从用户需求到行业实践，饱含技术干货和真知灼见的嘉年华大会吸引了上千人到场参与，上万人观看线上直播。

- **2023 可信数据库发展大会**

大会时间：2023.7.4

主办方：中国通信标准化协会中国信息通信研究院

主要内容：本届大会以“自主·创新·引领”为主题，共设置9个论坛，除7月4日主论坛外，7月5日分设金融行业、电信行业、互联网行业、汽车行业、云原生与开源数据库、搜索与分析型数据库、数据库运维及生态工具、时序时空及图数据库8个分论坛。近百位行业协会领导、数据库学术大咖、产业链各环节数据库负责人、资深技术专家将齐聚本届大会，带来极为丰富的主题演讲内容，与将要到场的1000+位开发者及关注数据库发展的行业人员，共同论道我国数据库高水平自立自强之路。

- **第十四届中国数据库技术大会（DTCC2023）**

大会时间：2023.8.16-8.18

主办方：IT168、ITPUB、ChinaUnix 技术社区

主要内容: 本届大会以“自主 · 创新 · 引领”为主题，共设置 9 个论坛，除 7 月 4 日主论坛外，7 月 5 日分设金融行业、电信行业、互联网行业、汽车行业、云原生与开源数据库、搜索与分析型数据库、数据库运维及生态工具、时序时空及图数据库 8 个分论坛。近百位行业协会领导、数据库学术大咖、产业链各环节数据库负责人、资深技术专家将齐聚本届大会，带来极为丰富的主题演讲内容，与将要到场的 1000+ 位开发者及关注数据库发展的行业人员，共同论道我国数据库高水平自立自强之路。

● **第 40 届 CCF 中国数据库学术会议（NDBC 2023）**

大会时间: 2023.10.13-10.15

主办方: 中国计算机学会

主要内容: 本届大会主要关注数据库领域所面临的新挑战、新问题和新方向，着力反映我国数据库技术研究的最新进展，为科研院所、科技企业的数据库研究、开发和应用相关人员搭建交流平台。大会共收到正式论文 156 篇，录用 70 篇，审稿人共计 207 名。大会由 5 个大会特邀报告、研究生学术辅导班、企业之夜、5 个分论坛报告、4 个大会论文分组报告、系统演示、软件学报专刊等一系列活动组成，并邀请了国内外数据库领域著名专家到会作专题报告，来自国内各大高校、科研机构、公司企业的 830 余人注册会议，线下参会人员超 900 人，会议期间，报告嘉宾和参会代表热情沟通、诚挚交流。

二、厂商及社区年度大会

除了数据库行业大会之外，众多厂商亦举办了 2023 年度大会，聚焦产品发布、联合解决方案、奖项等议题。还有例如行业趋势分析、技术交流、用户交流等。这些内容可以帮助厂商与用户建立良好的沟通，了解用户需求，促进数据库行业的发展。

表 9 2023 中国数据库厂商及社区年度大会一览

大会时间	主办方	大会名称	主要内容
2023.2.17-2.19	中国开源软件推进联盟 PostgreSQL 分会	中国 PostgreSQL 数据库生态大会	本届大会由中国开源软件推进联盟 PostgreSQL 分会&中科院软件所联合举办。现场百余位专家学者、厂商和用户代表，以及 CSDN 线上会场数千名 PostgreSQL 使用者和爱好者，共同见证此次盛会的召开。

2023.3.24	阿里云	瑶池数据库峰会	本届峰会集结国内外产学研大咖，凝聚知名数据库技术领袖、资深行业专家、生态合作伙伴和媒体同仁，共论云原生一站式数据管理与服务的发展之路。
2023.3.25	OceanBase	OceanBase 开发者大会 · 2023	3月25日，首届 OceanBase 开发者大会在北京举行。大会发布了 OceanBase 4.1 版本，公布两大友好工具，升级文档易用性，统一企业版和社区版代码分支，全面呈现了 OceanBase 打造极致的开发者友好数据库的成果。
2023.5.25	openGauss 社区	openGauss 开发者大会 2023	本次大会演示了 openGauss 最新版本技术特性，分享基于 openGauss 的行业联合创新成果及商业实践，展示社区最新治理进展，并讨论社区版本规划。在社区技术创新方面，会上发布了 openGauss 5.0.0 版本特性解读及 LiveDemo 展示；在社区优秀实践方面，会上由来自中国移动、联通云、邮储银行等伙伴进行了商业实践分享及成果展示；在社区生态成果方面，会上发布了基于 openGauss 的数据库一体机、伙伴能力认证和社区优秀开发者颁奖典礼。
2023.7.13	PingCAP	创新涌动于先—PingCAP 用户峰会 2023	2023年7月13日，企业级开源分布式数据库厂商 PingCAP 在京成功举办 PingCAP 用户峰会 2023。本届峰会以“创新涌动于先”为主题，PingCAP 全面解析了 AI 时代 TiDB 的演进方向，宣布 TiDB Serverless 正式商用。会上，PingCAP 携手用户代表发布平凯数据库，以更加完善的国产化生态兼容和企业级服务支持能力降低中国企业升级数据基础设施的成本和复杂性。30 多位来自各行业、深耕数据库领域多年的意见领袖分享了在技术涌现的时代，从技术领先到商业成功的创新故事。

2023.10.21	飞轮科技	Doris Summit Asia 2023: 与创新者同行	本届峰会以「与创新者同行」为主题，设置主论坛和智慧金融与政企、先进智造与电信、企业服务与新经济、互联网与文娱4个平行论坛，来自美团、银联商务、中国邮政储蓄银行、阿里云、腾讯、小米、华为、网易、平安、众安保险、无锡锡商银行、中国联通、中国电信、爱玛、雨润、奇安信、虎牙、货拉拉、趣丸科技、360、浩瀚深度、地平线、观测云、中铝视拓等行业创新引领者的数十位技术专家出席峰会并贡献了精彩的技术演讲，分享基于Apache Doris的行业最佳实践与多场景解决方案，共同探讨数据分析领域最前沿技术与未来趋势。
2023.11.03 -11.05	中国开源软件 推进联盟 PostgreSQL 分会	中国 PostgreSQL 数 据库生态大会	「2023 中国 PostgreSQL 数据库生态大会」于11月3-5日在北京中国科学院软件所召开！3天会议共邀请了30余位国内知名的社群专家，举行了主论坛、技术应用论坛、内核与新特性论坛、运维技能实践论坛，围绕不同的技术要点展开深入解读。除此之外，在最后一天开设了全天的私享培训日，进行体系化技能培训。
2023.11.17	镜舟科技	StarRocks Summit 2023	本次峰会以「极速进化，融合“新”生」为主题，40余场分享演讲在全天密集开展，来自平安银行、华润、腾讯游戏、阿里云、伊利、美的、京东等头部企业的大数据专家围绕数据进化、需求进化、技术进化，细致讲解大数据分析的最新技术和最佳实践，为大数据分析行业呈现了一场精彩的技术盛宴。
2023.11.26	OceanBase	砥砺前行，履践 致远-2023 OceanBase 年 度发布会	11月16日，2023 OceanBase 年度发布会在京顺利召开。在本次发布会上，OceanBase 对外正式宣布：将持续践行“一体化”新战略，为关键业务负载打造一体化数据库。同时，会上发布一体化数据库的首个长期支持版 OceanBase 4.2.1 LTS，标志着 OceanBase 一体化数据库进入可规模化上线使用的长期支持阶段。
2024.1.4	ACMUG 技术 社区	中国 MySQL 生 态年会	本次大会由中国 MySQL 用户组组织，集结了NineData、阿里云、PingCAP、万里数据库、美团、泽拓科技、Databend、华为云、沃趣科技等多位技术大咖及专家共同探索 MySQL 应用实践及未来展望。

12、数据库产业图谱

2023 年中国数据库产品数量达到 280 余款，覆盖图数据库、向量数据库、数据仓库、时序数据库等多种类型。墨天轮整理发布了全球图数据库产业图谱、全球向量数据库产业图谱、全球数据仓库产业图谱、全球时序数据库产业图谱、全球搜索数据库产业图谱、全球实时数据库产业图谱、全球时空数据库产业图谱、全球 PG 系数据库产业图谱、中国基于 MySQL 开发的数据库产业图谱等多个数据库分类图谱，以便于广大用户更清晰地了解和选择合适的数据库产品。

拓展：数据库产业图谱合集：<https://www.modb.pro/topic/639186>

全球图数据库产业图谱



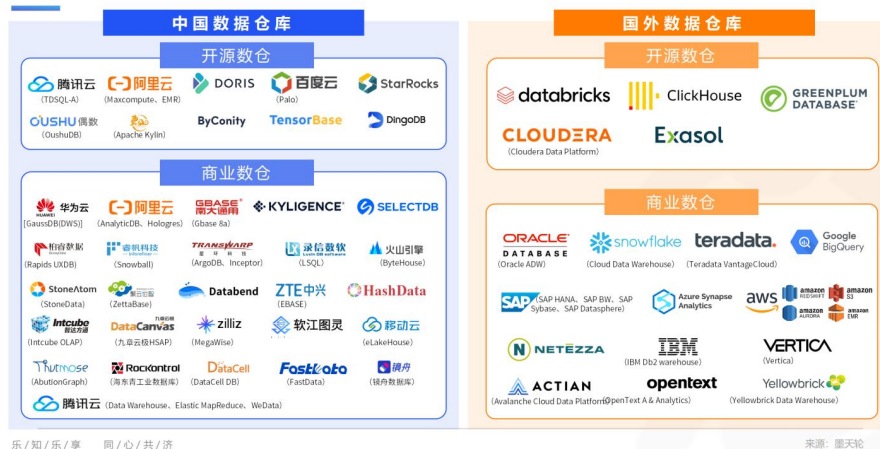
全球向量数据库产业图谱



全球数据仓库产业图谱

全球数据仓库产业图谱

墨天轮



全球时序数据库产业图谱

全球时序数据库产业图谱

墨天轮



全球搜索数据库产业图谱

全球搜索引擎数据库产业图谱

墨天轮



全球实时数据库产业图谱

全球实时数据库产业图谱



乐 / 知 / 乐 / 享 同 / 心 / 共 / 济

来源：墨天轮

全球时空数据库产业图谱

全球时空数据库产业图谱



乐 / 知 / 乐 / 享 同 / 心 / 共 / 济

来源：墨天轮

全球 PG 系数据库产业图谱

全球PG系数据库产业图谱



乐 / 知 / 乐 / 享 同 / 心 / 共 / 济

中国基于 MySQL 开发的数据库产业图谱

中国基于MySQL开发的数据库



乐 / 知 / 乐 / 享 同 / 心 / 共 / 济

来源：墨天轮

13、数据库社区年度发展概况

2023 年，中国数据库社区发展迎来了繁荣发展，多家数据库社区用户数不断攀升。越来越多的开发者、企业加入了社区贡献与生态建设中，为社区与产品发展添砖加瓦。作为数据库从业者与爱好者学习交流的重要平台，各用户社区组织发起了多种类型丰富的线上、线下交流活动，并与各大高校建立合作组织学习培训项目。当前，各社区活跃度均迎来了进一步提升，尤其是国产数据库自建社区活跃度日渐攀升，整体发展渐趋成熟。受篇幅所限，报告仅汇总介绍部分 2023 年较为活跃的数据库相关社区发展情况。

一、国外数据库的本土化用户社区

PostgreSQL 中文社区、中国 PostgreSQL 分会、ACMUG 社区等作为国外数据库产品的中国本土化用户社区，大多成立时间较早，其中部分社区与官方建立合作招募用户、部分社区均由民间志愿者组成，崇尚技术布道，均已建立了相对完善的运行制度及用户架构体系，会不定期组织发起技术沙龙、培训认证等活动以加强技术交流与生态共建。

● PostgreSQL 中文社区

PostgreSQL 中文社区是一个非营利的民间组织，自 2013 年 PostgreSQL 中国用户协会成立至今已逾 10 年，发展至今已帮助到了众多用户学习掌握专业技术，同时也帮助企业解决了人才培养和商用数据库成本问题，大力推动了 PG 技术在中国的发展。截至 2023 年，社区总注册用户数已达 13233 人，全年于杭州、西安、温州、武汉等

地举办了 6 场技术交流活动，多位社区 PG 专家与开源技术爱好者、资深开发者们线下相聚。



● 中国开源软件推进联盟 PostgreSQL 分会

中国开源软件推进联盟 PostgreSQL 分会（简称中国 PG 分会）于 2017 年成立，是开源联盟旗下致力于推动 PostgreSQL 技术应用与生态发展的组织，持续组织 PG 技术布道、发起培训与认证等。2023 年，社区持续发布技术干货、企业招聘，组织线下峰会沙龙、线下直播，在其 11 月举办的 2023 中国 PostgreSQL 数据库生态大会中，除汇集产学研等方面的学者、专家进行演讲分享、布道交流外，也为 2023 年度在 PG 生态表现优秀、贡献突出的专家、企业进行颁奖。



- 中国 MySQL 用户组

中国 MySQL 用户组（China MySQL User Group）简称 ACMUG，是覆盖中国 MySQL 技术爱好者的技术社区，受到 Oracle User Group Community 和 MariaDB Foundation 共同认可，于 2009 年正式成立，遵循“开源，开放，开心”的口号，交流 MySQL 使用经验、推广开源技术。2023 年，于上海、深圳、北京、成都、西安举办了五场线下活动，并于厦门成功举办中国 MySQL 生态年会。其中，MariaDB 创始人、MySQL 之父 Monty 曾出席【深圳 MariaDB 专场】活动。



- Elasticsearch 中文社区

Elasticsearch 中文社区成立已逾 12 年，曾自行组织编辑 Elasticsearch 权威指南中文版这一爱好者的经典入门教程。于 2023 年 11 月进行全新的品牌升级，正式更名为“搜索客”社区，并持续整理分发相关行业新闻的社区日报、周报，分享技术干货与沙龙信息。



- MongoDB 中文社区

MongoDB 中文社区成立于 2014 年，是大中华区获得 MongoDB 官方认可的中文社区，由博客、线下活动、技术问答、微信/QQ 群、官方文档翻译、源码及工具等版块组成。当前已搭建了完整的技术学习文档教程，并不定期分享 MongoDB 技术干货。

2023 年，社区先后组织举办了数据栈最佳实践、关系迁移器等主题的线上线下分享活动，并积极联动企业组织翻译 MongoDB 官方文档。

二、国产数据库开发者与用户社区

近年，随着国产数据库行业的不断繁荣发展，部分国产数据库选择开源发展，或建立自己的用户社区，以此加强用户技术交流、推动生态共建。2023 年，这些国产数据库社区进行了不断地探索与实践，通过组织各项活动与技术布道，增强核心技术实力的同时，加强了产品落地实践，同时增加了社区用户活跃度、打造了良好的自有技术生态。

● OceanBase 开源社区

2021 年 6 月 1 日，蚂蚁集团自研数据库 OceanBase 宣布开源，同时 OceanBase 开源社区正式成立，至今注册用户已达 75000+人。2023 年，OceanBase 社区专栏文章数超 920 篇、论坛答疑数达 39.6k；此外发起 43 场技术峰会、30 场城市开发者聚会、7 场企业交流、9 场线上直播分享等多场活动，促进了多行业 OceanBase 开发者互动交流、加强联络，线上线下打通大幅增加社区活跃度。



● PolarDB 开源社区

PolarDB 开源社区是阿里云数据库开源产品 PolarDB 的技术交流平台，自 PolarDB 已正式开源近 3 年来，秉持“打造世界级云原生数据库开源社区”的信念，已建立了 15 个 SIG 组，吸引 30000+贡献者和社区用户。通过兼容生态、携手生态伙伴，PolarDB 在不断推动开源协作与人才发展，吸引了韵达、网易数帆、龙蜥等 60 多家生态伙伴，10+所高校挂牌成立 PolarDB 开源数据库工作室，并与武汉大学、华东师范大学联合开展教学课程、智慧问答助手等多个项目。

● openGauss 开源社区

自 2020 年 7 月 1 日 openGauss 开源社区正式成立三年多以来,社区企业和学术机构数超过 600 家、开源贡献者逾 6100 人,版本下载量超过 230 万次,全量代码行数超 2100 万行。2023 年,联合合作伙伴举办线下 Meetup19 次、线上直播 120+期、大型峰会 5 次。据 Gitee 平台指数统计,openGauss 已经成为国内最活跃的开源数据库根社区。同时,openGauss 联合全国高校推出“基础软件百校种子计划”,已与 300 多家高校联合开设 openGauss 课程,培养 500 多位优质师资,40000 余位技术人才。



● TiDB 社区

2023 年, TiDB 开发者社区 GitHub 累计 Star 数达 37081、贡献者 2200+、服务国家和地区 47 个、已合并 Pull Request 26000+、已解决 Issue13,000+。此外 TiDB 用户社区用户数持续增长,截至 2024 年 1 月 1 日已达 33926 人,问答论坛积累优质技术问题 23000+,其中 90%的问题都得到了解决,累计优质回复达 249000+条,社区优质文章达 1100+篇。此外,2023 年 TiDB 社区共计发起 9 场 Meetup 活动及用户峰会等线下活动,组织当地社区用户畅谈数据库行业的技术趋势、DBA 职业发展。



● TDengine 开源社区

2023 年，TDengine 在 GitHub 上的 Star 数突破至 22.4k，开发者贡献的 Issue 数和 PR 数分别达到了 4416 和 19665，社区活跃度可见一斑。此外，在活动布道方面，为了更好地赋能开发者、助力行业数字化转型，2023 年 TDengine 共参加了 16 场颇具规模和影响力的行业及开发者峰会，其中包含国外 KubeCon + CloudNativeCon Europe 2023、ODSC East、Percona Live 2023 等知名盛会。

● StarRocks 社区

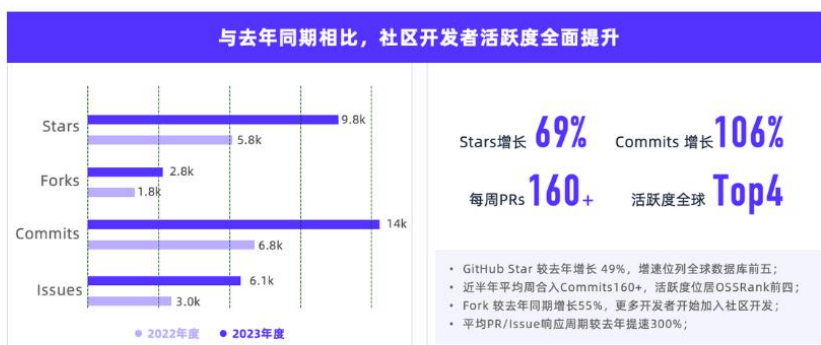
2023 年，StarRocks 社区在产品升级、用户规模、开发者规模的各个方面都取得了显著的进展，相继发布了 4 个大版本升级、50+ 个小版本迭代，350+ 位核心开发者加入社区贡献者行列，GitHub Star 数达 7000+ 个。此外，社区发起及参与了 61 场线上线下技术交流活动，地域涵盖深圳、成都、杭州、上海、新加坡等，主题涉及前沿技术解析、特性解读、应用实践等。



● Apache Doris 社区

Apache Doris 社区开发者规模和开发者活跃度在 2023 年均有增长，在 GitHub 上的 Star 数达 9800 个，与去年同期相比增长近 70%；贡献者规模已经增长至近 580 位，平均每月的活跃贡献者已稳定至 120 人左右。同时社区也建立了更加成熟稳定的 CR 流水线，每个合入的代码都会经过 3000+ 的测试用例，这也使得社区以极快速度迭代的同时，稳定性也得以保证。当前 Apache Doris 社区已聚集超过 30000 名数据库以及大数据相关领域的工程师，在全球范围的用户规模已经超过了 4000 家。

更加繁荣的社区生态，更大的开发者规模和更高的开发者活跃度



@稀土掘金技术社区

● Apache IoTDB 开源社区

目前，Apache IoTDB 开源社区已拥有超 260 位贡献者、Star 数超 4.1k，2023 年新增 85 万行代码。2023 年有详细的社区文档与社区内容产出，并持续在内容、活动、品牌运营方向发力，共计发布 130 篇技术解析、应用案例等内容，举办或参与活动 20+场，于北京、上海、无锡、南京等地与超 5000 名用户见面交流。



● GreatSQL 社区

自 2022 年 8 月 GreatSQL 社区官网论坛上线至今，已迎来 3000+社区用户、收到 400+个问题讨论，2000 条互动回复。2023 年，GreatSQL 开源数据库项目成功捐赠开放原子开源基金会，成为基金会旗下孵化项目。此外，组织了 10 余个线上论坛互动活动，创建了社区茶话会、社区月会；自建、参加了 20 余场线上直播分享、线下技术沙龙。



● Milvus 社区

Milvus 作为国产开源向量数据库于 2019 年 10 月 15 日正式开源，当前 GitHub Star 数已超 25000，下载数超 1000w。2023 年，社区全年举行近百场线上线下活动 15 场，共计 30k+参与者；累计发布技术视频 25 个、播放量达 47k+；此外，作为众多开发者交流分享的重要阵地，Milvus 技术交流群的规模也突破了 5000 人。

● ByConity 社区

2023 年 5 月，字节跳动团队研发的云原生数据仓库 ByConity 正式宣布开源，截至 2023 年底，其 GitHub Star 数达 1409、PR 数达 595，累计 61 位 Contributor 共建者参与代码贡献。2023 年，ByConity 社区主办活动超过 10 场，其中每月举办一次线上 webinar，为社区的开发者们介绍最新技术特性；此外在北京和上海主办了 2 场线下活动，以此与当地开发者建立更好的链接。



三、数据库垂直生态社区

墨天轮数据社区、dbaplus 社群、ITPUB 等数据库垂直生态社区，通过分发、汇集大量高质量、有价值的内容并发起组织多种丰富的线上线下交流活动，当前已发展成为数据库从业者、爱好者学习交流的重要阵地。同时，伴随着社区自有体系的不断发展成熟，2023 年这些社区在与数据库产业上下游、政府组织、开发者们的联系上均有所加强，逐渐成了数据库生态发展中重要的一环，在促进数据库领域技术发展的同时推动产业持续发展与落地。

● 墨天轮数据社区

墨天轮社区是国内专业的数据库技术社区，当前已覆盖国内 40 万数据库相关从业人员，围绕数据人的学习成长提供一站式的全面服务，包含课程、直播、文档/文章阅览、在线运维等等，致力于建设一个有温度的技术社区。2023 年，围绕数据库流行度排行榜及国产数据库发展发布官方解读、专家邀稿及国内知名企业/专家专访共 30 余篇，发布月度行业分析报告 12 期，总阅读量达 20w+；官方及社区用户共新增原创文章 2.6w+篇，为广大数据库从业者及爱好者学习交流提供了丰富的参考资料。此外，社区联合 ACUD（中国数据库联盟）于深圳、杭州、成都、西安等地成功举办 4 场技术分享活动并于北京举办了 2023 年数据技术嘉年华大会，为当地数据库从业者、爱好者搭建本地交流空间，共计数千人进行互动交流与学习实践。



● dbaplus 社群

dbaplus 社群是围绕数据库、大数据、AIOps 的企业级专业社群，坚持每天推送精品原创文章、每周举办线上技术分享、每月举办线下技术沙龙等，受众达 20W+。社群

由 2016 年主办的 Gdevops 全球敏捷运维峰会于 2023 年 8 月迎来全面升级，正式更名为 XCOPS 智能运维管理人年会，并于 11 月在广州成功举办，聚焦运维及数据库两大领域前沿技术与应用实践进行分享交流。

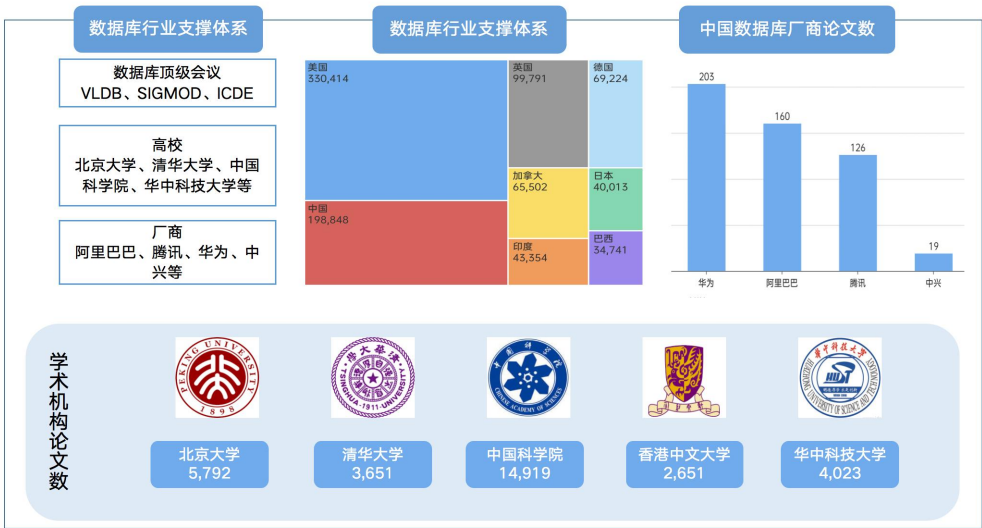
● ITPUB 社区

ITPUB 社区是 IT168 旗下的企业级业务落地平台，至今已有 22 年的历史，是中国权威的架构师、数据库、存储服务器领域的技术社区。2023 年，社区共计发起十余场云计算、数字化、数据库等线上主题交流活动，汇聚行业专家分享最新技术实践；并成功举办了第十六届中国系统架构师大会、第十四届中国数据库技术大会两场行业技术大会，后者汇集行业超百位专家重点围绕数据库内核、云原生数据库、图数据技术等内容展开分享和探讨，为大数据领域从业人士提供一场年度盛宴。

14、数据库高校学术力量蓬勃发展

一、数据库学术论文概况

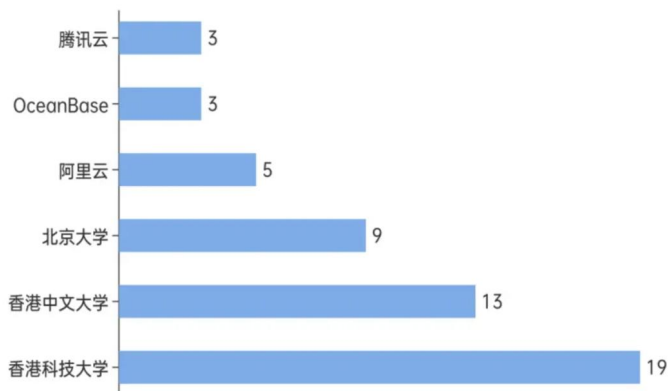
中国数据库领域的自主创新能力显著提升。2023 年中国在全球数据库领域的论文数量占比达 27.23%，位列全球第一。以下数据基于 Web of Science 核心合集，检索主题为 Database 的检索机构，检索日期为 2024 年 1 月。



2023 年数据库国际顶级学术会议 SIGMOD 仅收录 186 篇论文。其中由浙江大学与阿里云共同完成的《在数据库管理系统的连接优化器中检测逻辑漏洞》成果脱颖而出，斩获 2023 SIGMOD 最佳论文奖，实现了中国大陆研究团队在数据库国际顶会的历史性突破。

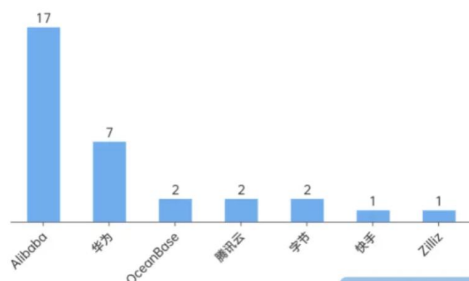


国内论文主要来源及数量

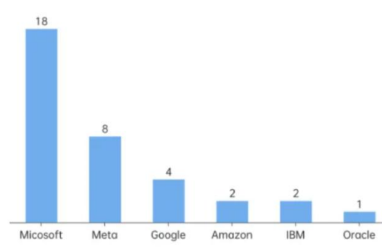


2023 年 8 月 VLDB 召开，VLDB2023 会议中共有 295 篇国内外论文入选。VLDB 会议全称 International Conference on Very Large Data Bases，是数据库领域历史悠久的三大顶级会议（SIGMOD、VLDB、ICDE）之一，每届会议集中展示了当前数据库研究的前沿方向、工业界的最新技术和各国的研发水平，吸引了全球顶级研究机构投稿。

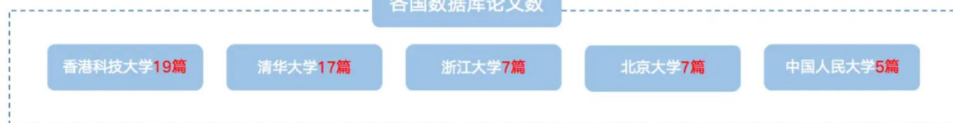
国内数据库厂商被收录论文数



国外数据库厂商被收录论文数



各国数据库论文数



二、数据库年度国赛

● 2023 年金砖国家职业技能大赛-信息技术应用创新赛项

2023 年 10 月 14 日-15 日，由南非高等教育与培训部、金砖国家工商理事会、南非豪登省政府主办，人大金仓协办的 2023 年金砖国家职业技能大赛-信息技术应用创新赛项华南赛区选拔赛成功举办。人大金仓在竞赛平台搭建、技术支持、专家指导等多方面积极支撑本次大赛，助力金砖国家技术技能人才培养。

● 鲲鹏应用创新大赛 2023

11月2日，以“数智未来因你而来”为主题的鲲鹏应用创新大赛2023全国总决赛在四川成都圆满落幕，经过长达6个月的层层筛选与激烈角逐，最终从3大赛事、5大赛道中评选出了13个金奖、16个银奖、19个铜奖。鲲鹏应用创新大赛是面向全球开发者的顶级赛事，旨在激发行业应用创新、加速产业融合、促进人才培养，吸引全产业开发者共同打造鲲鹏全栈解决方案。鲲鹏应用创新大赛自2020年创立以来已经连续举办了三届，今年第四届大赛在规模、赛题设置、赛队辅导等方面进行了全方位提升。在企业赛的基础上，增加了高校赛和科研赛，为高校团队提供了一个更加公平的竞赛环境，吸引了来自全国50多个顶尖院校的400多支队伍参赛，相比去年翻了两倍。



● 中国软件开源创新大赛

第六届“中国软件开源创新大赛”在国家自然科学基金委信息科学部的指导下，由中国计算机学会（CCF）主办，西北工业大学、绿色计算产业联盟、CCF 开源发展委员会联合承办。旨在为国内开源社区提供展示、交流、合作的平台，激发开源创新活力，培养开源实践人才，助力开源生态建设的高质量发展。

大赛面向国家“十四五”开源生态发展战略布局，聚焦“卡脖子”软件领域以及人工智能、大数据、芯片设计、物联网等前沿技术领域的开源软件，旨在为国内开源社区提供展示、交流、合作的平台，激发开源创新活力，培养开源实践人才，助力开源生态建设的高质量发展。

● 阿里云天池全球数据库大赛

2023年12月25日，第五届“天池数据库大赛”圆满收官，其是由阿里云主办，阿里云瑶池数据库、天池团队承办的数据库年度品牌赛事。自2018年以来，该比赛已连续成功举办了5届，吸引了来自国内外数千支优秀队伍和个人开发者参加。本届大赛采用双赛道机制，聚焦高性能共享存储、海量数据高效读写等核心业务场景设计赛题，吸引了全球7047支战队报名参赛。旨在以赛促学、以学促行，是集比赛、培训于一体的平台。



● 2023 全国大学生计算机系统能力大赛暨第三届 OceanBase 数据库大赛

2024 年 1 月 10 日 2023 全国大学生计算机系统能力大赛暨第三届 OceanBase 数据库大赛落下帷幕。作为专注于数据库内核技术的大赛，OceanBase 数据库大赛发起于 2021 年，旨在培养和发现数据库行业专业内核能力的关键人才，已连续举办三届，有超过 3500 支队伍参赛。2023 年，大赛进一步升级为全国大学生计算机系统能力大赛，由系统能力培养研究专家组发起，全国高等学校计算机教育研究会、系统能力培养研究项目示范高校共同主办、OceanBase 承办、华东师范大学等协办。



三、数据库高校生态合作

● 创邻科技与香港科技大学（广州）共建联合实验室

2023 年 3 月 30 日，创邻科技与香港科技大学（广州）合作共建的“创邻图数据联合实验室”在香港科技大学（广州）校园内正式揭牌。未来，双方将围绕万亿级大图神经网络计算框架、海量数据的时序图处理、分布式原生图数据库事务及性能优化等前沿图技术领域研究展开深入合作，充分发挥校企双方优势，提升国内技术研究水平，推动产学研协同创新发展。



- **浙大与阿里云合作成果斩获 2023 SIGMOD 最佳论文奖**

2023 年 6 月 21 日，数据库国际顶级学术会议 SIGMOD 在美国西雅图举行，由浙江大学与阿里云共同完成的《在数据库管理系统的连接优化器中检测逻辑漏洞》成果脱颖而出，斩获 2023 SIGMOD 最佳论文奖。浙大与阿里云研究团队提出了一种名为 TQS（转换查询合成）的新方案，通过引入机器学习等方法，创造性地解决了执行实现的正确性验证难题，以极小的计算代价自动探索更大检测空间，最终实现更完备的检测覆盖。



- **Oceanbase 与北理工大、新加坡国立大学联合研究成果被 SIGMOD 2023 收录**

2023 年 6 月 21 日，数据库国际顶级学术会议 SIGMOD 在美国西雅图举行，OceanBase 与北京理工大学、新加坡国立大学的最新联合研究成果《FEAST: A Communication-efficient Federated Feature Selection Framework for Relational Data》，被 SIGMOD 2023 Research Paper（研究类论文）收录。

- **PingCAP 协办中国数据库暑期学校**

第 22 届中国数据库暑期学校于 2023 年 7 月 21 日 - 7 月 26 日在福建龙岩举办，由龙岩学院数学与信息工程学院承办，中国人民大学和平凯星辰（北京）科技有限公司协办。中国数据库暑期学校下属于中国计算机学会数据库专业委员会。旨在为中国从事或有志于从事数据库理论与技术研究的研究生和青年学者提供一个学习和交流的机会。学院每年不定期地邀请数据库及相关领域内国际知名的学者来中国讲学，以促进我国对国际数据库学科前沿全面及时了解，并在此前提下，立足应用研发具有特色的数据管理技术和系统。

● 科蓝与清华大学成立联合研究院

2023年12月18日，清华大学与科蓝软件签署《联合建立清华大学-北京科蓝软件系统股份有限公司先进智能数据库合作协议书》，成立联合研究院。基于清华大学和科蓝软件在数据库和数字化方面的自主知识产权的产品基础和国内领先国际一流的技术与行业积累，在先进数据库领域开展长期稳定深度的合作，研发关键核心技术完全自主可控的高效、安全、智能的新一代数据库系统软件，实现在国家关键信息基础设施和重点行业的核心业务应用。



15、墨天轮 2023 年度数据库获奖名单

随着数字化转型深入推进和数据量的爆炸式增长，千行百业应用对数据库的需求变化推动数据库技术加速创新，全球数据库产业快速发展，我国已迈入第一梯队。2023年国产数据库在技术创新、市场竞争和国际合作等方面取得了显著的成就，展现出振奋人心的发展态势。

墨天轮以近 50 个客观中立的指标为基础，经过精心遴选，评选出了 2023 年的中国数据库行业重要奖项。我们期望通过这些数据展示中国数据库产业的全貌，为行业的持续发展提供支持与动力！接下来，让我们一同探索 2023 年度哪些中国数据库脱颖而出，成为行业的亮点。

一、2023 年度最具影响力数据库奖

“2023 年度最具影响力数据库奖”评选标准从技术能力、市场表现以及生态建设等三个维度出发，通过考量技术创新与研发能力、市场份额、财务状况、融资能力，以及顶级会议论文数、高校合作、专利数等近多项指标遴选出获奖产品，该奖项代表了产品作为国产数据库领军者，具备强大的品牌影响力以及行业凝聚力。

● OceanBase



产品介绍

OceanBase 完全自主研发，已连续 11 年稳定支撑双 11，创新推出“三地五中心”城市级容灾新标准，是全球唯一在 TPC-C 和 TPC-H 测试上都刷新了世界纪录的原生分布式数据库。产品采用自研的一体化架构，兼顾分布式架构的扩展性与集中式架构的性能优势，用一套引擎同时支持 TP 和 AP 的混合负载，具有数据强一致、高可用、高性能、在线扩展、高度兼容 Oracle/MySQL、对应用透明、高性价比等特点。

获奖理由

OceanBase 在“墨天轮中国数据库流行度排行”中连续 14 个月排名第一（截至 2024 年 1 月），并持续保持领先趋势。2023 年 OceanBase 持续做技术投入，**持续践行“一体化”产品战略**，实现单机分布式一体化、TP 和 AP 一体化、云上云下一体化。通过一个数据库、一套架构、一个技术栈、一份数据、一个引擎，构建支持多模型、多兼容模式、多租户、多工作负载、单机分布式、多基础设施的一体化数据库，用一个数据库满足 80% 的数据库场景需求。目前 **OceanBase 数据库已服务超过 1000 家行业客户**，**客户数年增长 150%**，其中 **30% 客户将其应用于核心系统**，OceanBase 正在成为核心系统升级的首选数据库。在金融领域，OceanBase 已成为市场占有率第一的分布式数据库。

- PolarDB



产品介绍

PolarDB 是阿里云瑶池数据库旗下、国内首款自研的云原生数据库，在存储计算分离架构下，PolarDB 利用软硬件结合等技术优势，为用户提供秒级弹性、高性能、海量存储、安全可靠的数据库服务。PolarDB 产品具有多主多写、多活容灾、HTAP 等特性，交易性能最高可达开源数据库的 6 倍，分析性能最高可达开源数据库的 400 倍，TCO 低于自建数据库 50%。100%兼容 MySQL 和 PostgreSQL 生态，高度兼容 Oracle 语法。PolarDB 同时支持分布式扩展，用户可根据业务需求选择合适的版本。PolarDB 分布式版支持分布式事务、全局二级索引等特性，让开发者像使用单机数据库一样、获得高性能分布式数据库的一体化体验。

获奖理由

以 PolarDB 为代表的云数据库向云原生纵深发展。据了解，PolarDB 在过去 3 年实现 400%的增速，目前用户数已超过 10000 家，广泛落地于政务、金融、电信、物流、互联网等领域的核心业务系统。云原生数据库 PolarDB 始终以客户需求为导向，已上线的云原生 HTAP、Serverless、PolarDB4AI、集中与分布式一体化等业内领先的新特性，帮助客户解决了众多业务问题。1 月 17 日，首届阿里云 PolarDB 开发者大会在京举办，会上，云原生数据库 PolarDB 发布“三层分离”全新版本，基于智能决策实现查询性能 10 倍提升、节省 50%成本。

- openGauss



产品介绍

openGauss 是一款开源关系型数据库管理系统，采用木兰宽松许可证 v2 发行。openGauss 内核深度融合华为在数据库领域多年的经验，结合企业级场景需求，持续构建竞争力特性。同时 openGauss 也是一个开源的数据库平台，鼓励社区贡献、合作。

获奖理由

openGauss 开源三年多以来，厚积薄发，一路披荆斩棘，不断积累技术经验，创新架构，拥抱开源，与产业，行业，开发者共建、共享、共治，在产业、生态、技术、商业四个方向上取得了巨大的进步，并在业界树立了标杆。三年期间，openGauss 社区企业数增长 100 倍，开源贡献者增长 50 倍，版本下载量增加 38 倍，代码量增长 16 倍，从 2020 年 6 月开源的 130 万行代码发展到今已经 2100 万行，规模应用于十大关键行业核心场景，2023 年线下集中式新增市场份额达到 21.9%，跨越生态拐点，从拓展期进入快速发展期。此外，openGauss 已有 17 家发行版伙伴，8 家 OGSP 伙伴，4 家一体机伙伴，有效支持了千行万业的数字化转型。从拓展期到发展期，openGauss 就像一匹黑马，在数据库领域掀起一股热潮。

二、2023 年度卓越表现数据库奖

“2023 年度卓越表现数据库奖”旨在表彰在国内乃至国际上占有一定市场份额，并在关键行业和领域得到了广泛应用的数据库产品。除市场份额与占有率外，还综合考量了技术实力、产品创新、应用领域、投产数量、业务覆盖率等多维度指标，代表着获奖产品在行业中具备深厚的市场影响力。

● 达梦数据



产品介绍

达梦数据是国内领先的数据库产品开发服务商，公司向大中型公司、企事业单位、党政机关提供各类数据库软件及集群软件、云计算与大数据产品、数据库一体机等一系列数据库产品及相关技术服务。

获奖理由

达梦数据是国内领先的数据库产品开发服务商，是国内数据库基础软件产业发展的关键推动者。多年来，公司始终坚持原始创新、独立研发的技术路线，公司核心团队在数据库领域拥有 40 余年研发经验及技术积累。目前，公司已**掌握数据管理与数据分析领域的核心前沿技术，并拥有主要产品全部核心源代码的自主知识产权**。产品成功应用于金融、能源、航空、通信、党政机关等数十个领域。根据赛迪顾问及 IDC 发布的报告显示，近年来公司产品市占率位居中国数据库管理系统市场国内数据库厂商前列。

● TDSQL



产品介绍

TDSQL 是腾讯云自主研发的一款金融级分布式数据库产品，旗下涵盖金融级分布式、云原生等多引擎融合的完整数据库产品，提供业界领先的金融级高可用、存算分离、企业级安全等能力，同时具备智能运维平台、Serverless 版本等完善的产品服务体系，是金融及众多企业客户核心替换的优选数据库。

获奖理由

腾讯云数据库 TDSQL 已应用于超过 50 万客户，TDSQL 已经被 4000 多家来自金融、公共服务和电信等垂直行业客户采用，服务超过 30 家金融机构完成核心系统替换，中国十大银行中的七家都应用了 TDSQL，在 TOP 20 银行中也服务过半，**头部商业银行中 90%客户采用腾讯云的方案，其中七成应用在核心或关键业务领域**；头部前十的券商全部选择腾讯云服务，其中私有云平台合作行业领先；十二大保险集团中，七家选择与腾讯云深度合作，在不同金融机构核心系统中的渗透率均有显著提升。

● GoldenDB



产品介绍

GoldenDB 是金融市场排名第一的金融级交易型分布式数据库, 历经 20+年技术积累, 引领国产数据库自主创新, 专利超过 500 项, 实现内核 100%自主掌控; 成功通过金融、运营商行业现网多年严苛考验, 服务超百家重点行业客户, 引领各行业核心业务数据库分布式架构转型; 在行业主管部门的指导下引领标准制定, 与行业伙伴共建新生态, 打造国产数据库第一品牌。

获奖理由

作为金融级交易型分布式数据库, GoldenDB 支撑大型商业银行核心业务稳定运行超 4 年, 并支撑 50+客户核心业务、关键业务, 实现上百家客户的合作。GoldenDB 在金融领域已形成领先落地深度与广度, **银行业金融级分布式数据库市场份额占比 24.4%, 银行核心系统（不包括次核心）市场投产数量占行业 50%, 次核心及非银核心系统市场投产数量占行业 32%, 三项数据均为行业第一。**GoldenDB 以第一份额中标中国移动(70%)和中国联通(60%)分布式数据库集采, 现已在中国移动 3 大基线省份完成账务计费核心系统数据库投产, 在 20 多个省份进行核心数据库改造。GoldenDB 在核心业务覆盖的深度与广度为行业领先, 落地成果为行业提供丰富的参考案例, 对行业发展具有积极影响力。

三、2023 年度最具潜力数据库奖

“2023 年度最具潜力数据库奖” 评选标准定位于年度排名第 11 至第 20 名之间, 并在墨天轮排行榜年度得分平均值、创新价值、市场拓展计划、技术升级以及未来发展战略等方面呈现出强大潜力的数据库产品。

● YashanDB



产品介绍

YashanDB 是深圳计算科学研究院基于原创理论自主研发,内核代码自主率 100% 的新型数据库系统,提供包含单机/主备、共享集群、分布式实时数仓、空间数据库的 4 大产品形态,覆盖 OLTP/HTAP/OLAP 场景,一站式满足关键行业核心系统的多元应用需求。

获奖理由

YashanDB 自 2022 年正式推出以来,凭借其卓越的技术实力和市场潜力,在墨天轮上的评分和排名持续攀升。在技术层面,YashanDB 多篇数据库论文被 SIGMOD、ICDE 等国际顶会收录,数据库基础理论研究实力深厚;2023 年 YashanDB 突破了数据库领域技术“制高点”-共享集群技术,为关键核心业务实现国外数据库“1:1”平替。在市场拓展层面,YashanDB 客户群体不断扩大,已成功应用于政府、金融、能源、交通等重点行业数十个重要核心业务系统,并提供了可复制性的经验,为千行百业赋能。

● MogDB



产品介绍

MogDB 是云和恩墨基于 openGauss 开源内核进行增强提升,推出的一款安稳易用的企业级关系型数据库。云和恩墨结合自身十余年的技术沉淀和实践积累,投入一流内核研发团队,将商业数据库的能力带入 openGauss 社区,打造 MogDB 核心竞争力,实现高安全、高性能、高兼容、多场景支持等特性,同时在压缩技术、并行、自治、可观测性、一体机等方向重点创新,能够满足从核心交易到复杂计算的企业级业务需求,现已助力金融、电信、制造、政务等行业用户实现核心系统升级。

获奖理由

MogDB 在过去一年里在技术创新、能力表现、商业落地等方面都取得了显著成绩。MogDB 面向客户痛点并结合云和恩墨对成熟商业数据库的理解与技术洞察,围绕极致高可用、高性能密度、兼容能力增强、易用性提升等核心价值点推出大批创新特性,并

围绕实际应用场景配备整套工具链，方便用户进行新系统部署或者国产化替代。一年来，MogDB 顺利通过 CCRC EAL4+认证和 GB18030-2022 最高级（实现级别 3）认证，获评“可信数据库”，并继续扩大在金融、政务、高端制造等领域的落地应用，多个实践案例获荣誉奖项。MogDB 的后续发展值得期待，由此获评 2023 年度最具潜力数据库奖。

● SUNDB



产品介绍

科蓝软件 SUNDB 是国内稀缺的拥有完整知识产权与核心技术的原生分布式数据库，拥有四种产品形态：通用磁盘数据库、内存数据库、磁盘内存融合库、极速数据库；支持四种部署模式：单机版、集中式、分布式、云部署。SUNDB 数据库致力于解决国家信息安全问题，产品规划清晰，战略布局专注银行、保险、证券、电信、交通、能源、军工等国家关键信息基础设施的核心业务系统，拥有国家关基领域超大型生产系统多年稳定运行的典型应用案例。

获奖理由

科蓝软件 SUNDB 是一款分布式关系型数据库，立足金融级数据库高安全、高要求特性，是真正掌握根技术的中国数据库，可以向全球合法销售 License。2023 年底，科蓝软件联合清华大学成立“先进智能数据库联合研究院”，共同推动数据库核心技术创新研发。目前科蓝软件 SUNDB 已首批通过“可信数据库”评测，能满足金融级企业核心业务系统需求，并可在国内关键基础设施领域替代国外同类数据库产品。

四、2023 年度图数据库

“2023 年度图数据库” 评选标准是通过墨天轮排行榜排名、生态建设、专利数、顶会论文数、市场份额、入选国际权威报告等近 50 个综合指标遴选出来，奖项代表了该产品在图数据库中具有强大的号召力和凝聚力。

- TuGraph



产品介绍

高性能图数据库 TuGraph 由蚂蚁集团和清华大学共同研发，历经蚂蚁万亿级业务的实际场景锤炼，多次刷新国际图数据库基准性能测试 LDBC-SNB 世界纪录，在功能完整性、吞吐率、响应时间等技术指标均达到全球领先水平，在金融、能源、政务、社交网络等领域都有巨大应用潜力，为各行业企业管理和分析复杂关联数据提供了高效、易用、可靠的平台。

获奖理由

作为图数据库领域的佼佼者，TuGraph 在 2023 年取得了瞩目成绩。TuGraph 拥有业界领先的集群规模和性能，已应用于蚂蚁集团万亿规模的业务（例如支付宝），并支持超过 300 个业务应用（例如“双 11”），帮助金融风控效率显著提高，支付宝的资损率低于十亿分之一。2023 年，TuGraph 成为墨天轮中国图数据库流行度排行榜上摘取榜首次数最多的产品，并跻身 IDC 中国图数据库厂商“领导者”象限，在产品能力、技术战略维度评分上均处于头部水平。TuGraph 在 2023 年国际图数据库性能基准测试“LDBC SNB-BI”中打破世界纪录，在数据规模和性能方面领先约 2-3 倍。继开源单机版图数据库后，2023 年 TuGraph 先后开源了首个工业级流式图分析引擎和大规模图学习引擎。

五、2023 年度时序数据库

“2023 年度时序数据库”评选标准是通过墨天轮排行榜排名、生态建设、专利数、顶会论文数、市场份额、入选国际权威报告等近 50 个综合指标遴选出来，奖项代表了该产品在时序数据库中具有强大的号召力和凝聚力。

- KaiwuDB



产品介绍

KaiwuDB 是业内首款面向 AIoT 的分布式多模数据库产品，将自适应时序引擎、事务处理引擎、预测分析引擎、超速分析引擎、内存实时引擎等深度融合，可从容应对物联网时代多源异构数据的管理需求。基于“就地计算”核心技术，KaiwuDB 具备每秒百万级写入、毫秒级数据读取能力，满足海量、高并发的时序数据高速写入、快速查询、复杂查询等需求。

获奖理由

面向以海量时序数据为一大主要特征的 AIoT 领域，KaiwuDB 创新研发了“就地计算”等核心技术，实现了物联网场景下海量时序数据的高速入库与极速查询；提供高至百倍数据压缩能力及数据降采样存储，大幅降低成本；支持物联网场景下的流式计算、云边端协同，拥有数据库自治、高可用、高兼容、高安全等特性。2023 年，产品获得信通院“可信数据库”时序基础能力、性能、稳定性全项测试认可。目前，已在工业物联网、数字能源、智慧产业等快速发展的新兴数字产业形成可复制解决方案，打造了包括山东重工、国网山东综能在内的标杆应用项目。基于 KaiwuDB 的时序领域标杆案例分别获 2023 数博会“优秀科技成果”，信通院“数据库星河标杆案例”及 IDC “可持续发展先锋案例”。

总结：

回顾过去，我们看到数据库行业在面临巨大挑战的同时，也展现出强大的适应力和创新力。2024 年也将是数据库行业蓬勃发展的一年。我们对未来怀揣着信心，期待着更多技术创新的涌现，更多市场领域的拓展，以及更加健康繁荣的生态系统建设，推动着国产数据库技术不断迈向新的高度！

第二章：数据库关键技术及发展趋势

1、云原生

一、云原生数据库是什么？

云计算的出现促进了企业信息技术的进步。云计算供应商把大量的计算、存储和通信资源汇聚在一个“池”中，允许企业或个人用户根据需求购买这些资源，从而能快速、低成本地建立信息系统。当系统负载发生变化时，可以根据需要增减计算资源。对于供应商来说，能统一管理所有用户使用的大量资源，实现规模效应，降低边际成本。对于云计算用户来说，获取资源快捷、便捷，按需使用的成本更低。从全社会角度看，整体资源的使用效率得到提高，更环保。

数据库是最常用的基础软件之一，它能提供计算和存储能力。存储是数据的基本功能，计算能力则体现在数据库能完成用户的复杂分析和计算请求，以及进行查询优化、事务处理、索引维护等内部计算。

单机数据库部署在普通主机上，其存储和计算能力受硬件限制，难以扩展。分布式数据库可以通过增加机器来扩展容量和计算能力，但依然受到机器资源的限制。如果简单地将它们迁移到云上，将普通主机换成云主机，可能会出现一些问题，如网络瓶颈、写放大问题等；并且不能充分发挥云计算的优势，如资源弹性管理、按需使用，也无法灵活使用各种云存储产品。

如果分布式数据库如果未经改造，简单地迁移到云上；尽管分布式数据库有很好的扩展性，但不能算是“云原生”。首先，它弹性扩展的单位是“机器”，而不是更细粒度的“计算和存储资源”。其次，它在设计时并未考虑云平台的特性，没有进行相应的优化，以达到最优性能和成本。

“云原生数据库”的核心是设计一种更符合“资源弹性管理”理念的数据库架构，充分利用云平台的池化资源，适应云平台的基础设施。并且，经过多次迭代更新，产品已在云上发展成熟，云原生数据库的技术也普惠到了更多的场景和部署环境，包括私有云场景，为企业提供数据安全可控的解决方案。

二、云原生发展历史

云数据库的进化轨迹始于 2010 年代初，这个时期正值云计算技术的兴起，大量企业开始尝试将传统数据库搬迁至云端。近些年来，随着云基础设施的迅猛发展，云数据库也得到了飞速的拓展，并且因其按需扩展和按需付费等卓越特性，受到了中小企业和

互联网客户的广泛欢迎。然而，云数据库并非专为云场景或云环境设计和构建的，它仅仅利用了云的资源。因此，它存在一些固有的问题，如存储空间浪费、计算资源浪费、恢复时间目标大以及数据延迟、系统性能受限、网络带宽消耗大等，这些问题阻碍了业务的进一步发展。

Amazon 首先意识到上述问题，推出的云数据库 Aurora 就是为云计算时代而专门定制的一款关系型数据库。从此数据库又进入了一个崭新的阶段，云原生数据库。各大公有云厂商基于这一路线紧随其后，在优先保证上云兼容性的前提下，基于存算分离架构对传统数据库进行改造：通过把大量的日志操作放到后台异步处理，实现存储独立扩展，解决了 MySQL 数据库单库的数据量不能太大的痛点。并且云原生数据库提供了又兼容又能扩展的能力，在存储层面实现了扩展的同时，又保留了计算层面的不变和兼容，从而基本实现了完全的兼容性。像典型产品就有阿里云的 PolarDB，百度智能云的 GaiaDB 和腾讯云的 TDSQL-C。可以完美兼容传统的使用习惯，对交易类场景可以提供很低延迟的写事务能力。同时读扩展性与存储扩展性由于借助了分布式存储池化能力，也得到了极大增强。

而分布式数据库为了解决扩展性问题，是另外一种演进路线，则优先将系统的扩展性放在首位，通过规模来解决各类业务对写扩展能力的要求，像 OceanBase 和 TiDB 就是两个比较典型的产品。它们的典型特征就是将事物系统和锁系统角色拆分为单独的模块负责，计算层通过与这些模块交互实现了多个节点都可接受写请求，然后由统一的新事务+锁中心节点来进行仲裁。这样对于写负载本身需要较多计算资源的场景下会有很好的提升，但是由于事务和锁都需要跨网络进行交互，所以事务延迟是相对比较高的，在锁负载较重的负载下会成为一定的瓶颈。



云原生数据库由国外传入国内。如今,以阿里云的 PolarDB、百度智能云的 GaiaDB、腾讯云的 TDSQL-C 等为首的主要厂商,都在投入大量资源进行研发。短短三年的时间,市场已经形成了相对成熟的云原生数据库应用模式,并且已经在不同的场景中得到应用。可以看出,虽然国产云原生数据库的起步相对于国外稍晚,但其在国内的发展速度极快,影响力已经逐渐超过了国外的云原生数据库。

阿里云在 2024 年 1 月 17 日隆重召开 PolarDB 开发者大会,分享了阿里云 PolarDB 未来发展的四大趋势:深化服务化、多主多写、容灾、全球化等的云原生化、支持标准的 API、IAC、CI/CD 等工具的平台化、实现不同数据处理场景的一体化、与 AI 深度结合的智能化。

百度于 2023 年 12 月 20 日的智算大会上正式发布自研云原生数据库 GaiaDB 4.0,帮助解决企业复杂查询难题。GaiaDB 4.0 增强了并行查询能力,突破单机计算瓶颈,实现跨机多核并行查询,在混合负载和实时分析业务场景中性能提升超过 10 倍。针对不同的工作负载, GaiaDB4.0 推出列存索引和列存引擎,提升不同规模数据的查询速度,其中列存引擎最大可支持 PB 级数据的复杂分析,并且与事务处理业务严格复杂隔离。

总而言之,在互联网和云计算快速发展的时代背景下,各行各业对于数据库的需求在不断增加和变化,随着这些新的需求越来越广泛地被提出,用户意识到采用传统单一的数据库来应对各类场景的时代已经过去,各厂商提供云原生数据库需要从多方位实现资源规格的灵活控制、应用的多模、更优的弹性扩展能力、更好的成本控制方式等。

三、云原生数据库典型场景和核心优势

云原生数据库适用于多种应用场景,包括电商、金融、物流、社交媒体等。在这些场景中,云原生数据库能够提供性能强大、高可用性和可扩展性的数据服务,以满足各种业务需求。

(一) 承接高弹性流量

与大多数传统业务的稳定性和容易预测的资源需求相反,互联网业务可能会随时面临流量激增,这是传统线下部署方式难以处理的。云原生数据库凭借其出色的可扩展性,能根据业务负载灵活调整资源,实现实时扩缩容,用户无需担忧流量激增导致的存储资源不足。通过 Serverless 技术,云原生数据库可以实现最大程度的弹性扩展并按需付费。

云原生数据库（阿里云的 PolarDB、百度智能云的 GaiaDB、腾讯云的 TDSQL-C）都支持分钟级弹性扩容，能迅速应对如双 11 等大流量场景，相比传统本地业务模式，云原生其强大的扩缩容能力更具优势。

（二）保障高可用性

数据库采取多种方法确保业务的高可用性。云原生数据库基于计算存储分离架构，解决了主备切换的延迟问题，提升了性能并确保系统和业务的连续性。它不仅具有分布式数据库的高可用性，还能依赖云基础设施在单个可用区或整个地域宕机时，及时迁移和自动修复数据。通过自动全量增量备份和快速恢复，以及利用云存储的高可用性、异地容灾和快照能力，实现了更快、更可靠的备份恢复。

特别的，阿里云的 PolarDB 集群版和百度的 GaiaDB，分别通过全球数据库 GDN 功能和热活集群组三副本一致性存储，进一步增强了跨地域灾备和数据可靠性。

（三）满足高时效性

在当前高并发、多变的业务环境中，对数据库的时效性要求愈加严格。云原生数据库在运维方面，能在几分钟内完成安装，远快于传统数据库，且其版本升级的时效性也更高。在面对大规模业务负载时，如报表查询的高峰期，传统数据库需要大量硬件资源和预留服务器来保证查询效率，这不仅成本高，扩容速度也慢。而云原生数据库能快速扩容，满足业务高峰需求，有效支持业务快速变化。在性能上，云原生数据库采用计算存储分离架构（如发展历史中介绍），借助日志即数据、算子下推和并行计算技术，充分利用存储资源的计算能力，有效提升大型表扫描等网络数据流量大的业务场景的性能。

阿里云 PolarDB 支持秒级增加、修改和删除字段，可以大幅提升 SaaS 海量表维护的效率，大幅降低客户侧运维工作量。百度文库业务通过使用云原生数据库解决了传统数据库主从延迟高，导致数据同步形式在异常情况下会出现延迟，影响付费及文档服务体验的问题。

（四）应对混合型（AP+TP）业务

伴随企业的发展和业务的复杂化，我们经常看到大量不同类型的数据被存储在一起。以交互系统和报表系统为例，前者是 OLTP（在线事务处理）应用场景，后者则是 OLAP（在线分析处理）应用场景。如果这些数据被存储在同一个数据库中，数据库就需要在处理事务和分析数据时都表现出高效性。此外，有些用户可能会将多年的历史数据存储在一起，没有进行冷热分离。当查询这些历史数据时，由于访问量增大，可能会占用大量的 CPU 和内存资源，造成系统性能波动，影响在线事务处理业务。

阿里云的 PolarDB 和腾讯云的 TDSQL-C 通过支持列存储和向量化执行引擎，能在 OLAP 场景下显著提高查询性能。

（五）满足成本的诉求（弹性+智能化）

在传统数据库维护中，运行数据库实例涉及预备硬件和折旧费用，但主机资源常常未被充分利用，造成浪费。如果业务有高低峰，必须始终准备最高资源，使维护成本高昂。相比之下，云原生数据库降低了初始投入，且在业务增长时能够弹性扩展，按需使用资源，避免浪费。阿里云和腾讯云的 Serverless 集群成本可降低 80%。在存储和计算分离架构下，百度网盘迁移至百度云原生数据库 GaiaDB，资源利用率提高 50%，大幅降低成本。

另外，数据库运维和 SQL 优化等工作需由 DBA 完成，多数企业 DBA 资源有限，增加运维成本。云原生数据库的智能运维结合了 AI 算力和强大的数据库管控平台，具有高度自动化运维能力，易于管理数据库，缓解 DBA 压力，节省企业成本。

四、云原生数据库技术最新进展和发展趋势

（一）多级高可用，提供异地灾备能力

数据库作为底层数据存储环节，其可用性与可靠性直接影响系统整体。而线上情况是复杂多变的，机房里时时刻刻都可能有异常情况发生，小到单路电源故障，大到机房级网络异常，无时无刻不在给数据造成可用性隐患。

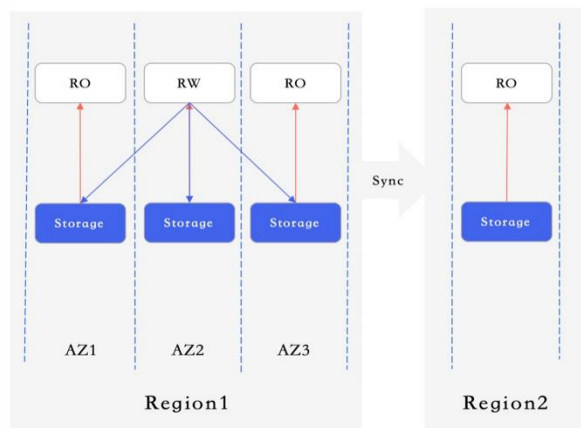
作为商业数据库，具备多级高可用能力是最核心的必备能力。这样才能抵御不同级别的异常情况，有力保障客户业务的平稳运行。阿里云提供全球数据库网络服务（Global Database Network，简称 GDN）是由分布在同一个国家内多个地域的多个 PolarDB 集群组成的网络。保障不论业务部署在一个或多个地域，都能通过全球数据库网络实现异地容灾。当主集群出现地域级别的故障时，您只需要手动将您的业务切换到从集群。百度智能云的 GaiaDB 支持多副本、跨可用区、跨地域三级别高可用，创新性地实现了多可用区热活高可用、单个实例支持跨可用区部署。在不增加成本的情况下，每个可用区均可提供在线服务，任何可用区故障都不会打破存储一致性。

数据库跨可用区高可用

- 使用物理同步协议，任意切换不丢失数据
- 多可用区热活同时提供服务，机房级故障无感
- 就近访问，低延迟平滑兼容慢节点
- 增强型自校验日志流，链式自校验，抵御未知错误

数据库跨地域高可用

- 跨地域物理复制，延迟仅有数十ms
- 支持跨地域快速切换
- 地域级就近读取



(二) 多级 HTAP 释放潜能

将内存池技术与 HTAP 结合是未来数据库技术的一个重要趋势。随着云原生数据库的普及和发展，OLTP 和 OLAP 能力的融合已经成了一个基本要求。在此基础上，结合内存池软硬协同技术，可以进一步实现网络吞吐的大幅度缩减，提高数据库的性能和响应速度。

阿里云的 PolarDB 目前已经支持内存池技术（内存资源进行统一管理和分配），并且可以有效地提高内存的利用率和系统的整体性能。将内存池技术与 HTAP 结合，可以利用内存池的高性能和低延迟特性，为 HTAP 提供更加稳定和高效的支持。

具体来说，通过将内存池技术与 HTAP 结合，可以实现以下几个方面的优化：

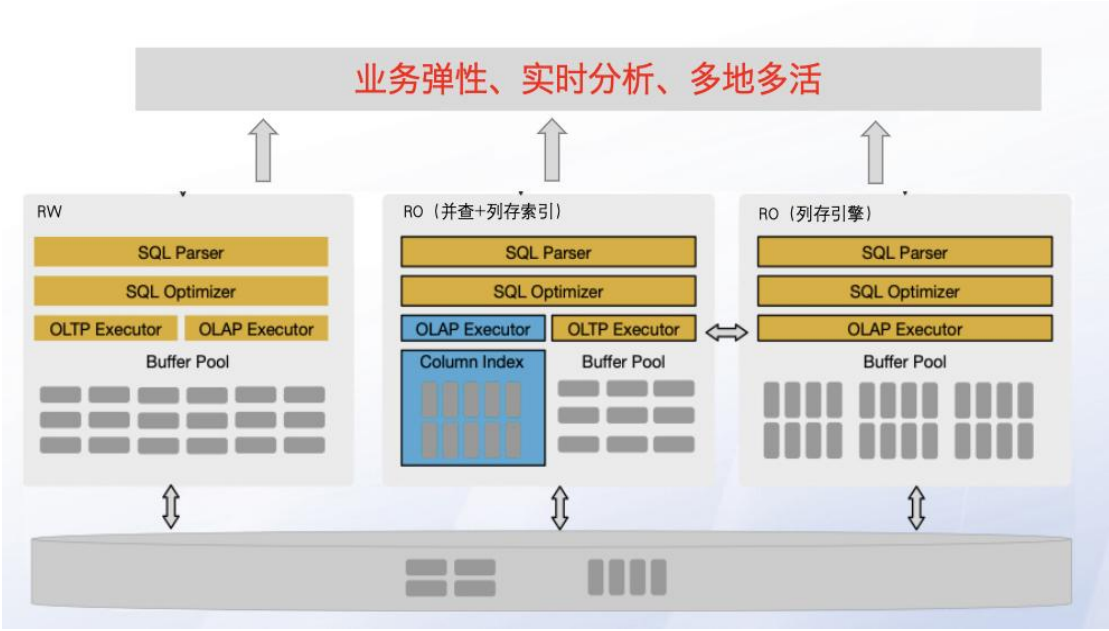
1. 减少网络通信开销：通过内存池技术，可以将数据缓存在本地内存中，减少了对远程存储的访问和网络通信的开销，从而提高了系统的整体性能和响应速度。
2. 提高数据一致性和可靠性：内存池技术可以保证数据的一致性和可靠性，避免了因为网络延迟或数据同步不及时导致的数据不一致问题。
3. 实现资源共享和动态扩展：通过内存池技术，可以实现资源的共享和动态扩展，使得系统可以根据实际需求进行灵活地扩展和收缩，提高了系统的可扩展性和可用性。

而百度智能云的 GaiaDB 实现了多级 HTAP，通过并行查询，内存索引，列存索引进行优化提升性能，具体如下：

1. 并行查询，传统 MySQL 数据库在单机的并行方面很弱，Gaiadb 并行查询在改善单机查询性能的同时，突破单机计算瓶颈，实现跨机多核并行查询，面向混合负载和实时分析业务场景，性能提升十倍以上。
2. 随着数据量越来越大，业务越来越多，列存是比较有效解决复杂查询和分析的办法，Gaiadb 针对不同的 workload 推出列存索引和列存引擎。兼容地解决了不同规模的数

据的查询加速的诉求。其中列存引擎最大可以支持 PB 级别的复杂分析,并且和 TP 的业务是严格复杂隔离的。TP 和 AP 的数据是通过 zeroETL 来实现透明的数据同步。

3. 为了进一步突破性能的天花板, Gaiadb 在内核数据流上进行深度优化, 包括共识协议优化: 单跳提交, 吞吐提升 40%, 时延降低 30%; 链路优化: 异步批量提交+自适应用户态 TCP, 时延降低 60%; 自适应动态回放存储多版本, 高负载读 IO 降低 50%; 通过这一系列优化让 gaiadb 整体性能, 成功做到整体性能大幅提升 60%。



(三) Serverless 已成为趋势，智能弹性助力降本增效

随着云计算技术的进步, Serverless 数据库成为发展趋势。然而, 仅提供无服务器计算无法满足所有用户需求。未来, Serverless 数据库需具备智能弹性能力。

目前, 阿里云和腾讯云已支持 Serverless 集群, 提供降低成本和提高效率的解决方案。未来, 需要发展智能弹性能力, 以满足用户需求, 提升系统性能和稳定性。智能弹性指数据库根据历史负载数据, 自动计算用户画像, 快速预测未来负载曲线, 并预先为弹性伸缩准备资源。这使得数据库能迅速响应需求变化, 避免负载达到资源上限, 减少资源浪费。智能弹性是 Serverless 数据库未来关键能力之一, 能帮助数据库处理大量数据时保持高效性能, 确保系统稳定, 降低运营成本, 提高资源利用率, 提升用户体验。

(四) AI 助力，构建智能化数据库

自 2023 年开始的 AI 时代。底层基础设施变成了 GPU 和 AI 能力。应用也变成了 AI 原生应用, 如海外比较火的 Jasper、Midjourney, 微软的 Copilot 等。在数据库行业我们看到至少两个方向, 一个是 AI4DB, 其中包括阿里的 DAS、百度的

DBSC 等，主要是利用 AI 大模型能力和专家经验实现数据库的智能化洞察、评估和优化。有效保证数据库服务的安全、稳定及高效。这意味着数据库将能够自动检测和诊断问题，进行自我调优和运维，同时保证数据的安全性和可靠性。另外一个方向就是 DB4AI，目前主要是向量数据库。向量数据库在解决大模型幻觉等方面，有非常不错的效果，是一个有潜力的细分赛道，头部公司估值已经达到 10 亿美元。

总的来说，未来的云原生数据库将继续与 AI 内外协作，包括 AI 帮助云原生数据库提效以及云原生数据库和向量能力的结合，向全场景智能数据库迈进，为各种应用场景提供更加智能、高效和安全的数据管理服务。

（五）云原生数据库变得更加普惠

云原生数据库目前正处于发展的攀升阶段，云厂商正在不断提升产品的性价比、性能以及稳定性。对大部分客户来说，他们更倾向于选择托管类的云数据库，主要因为这些数据库基于成熟的开源数据库构建，如 MySQL、Redis 等，这些数据库既成熟又普惠。企业认为这些完全可以通过技术进行控制，即便考虑到未来可能从云端迁移到自建 IDC，业务也不会与云厂商紧紧捆绑。

在选择云原生数据库时，企业也需要考虑迁移成本、技术可控性、云厂商的绑定以及总体拥有成本（TCO）等问题。目前，大部分的云原生数据库仍然依赖于云厂商自身的技术底座与硬件环境。

未来的云原生数据库应该朝向普惠化发展，一个真正普惠的数据库需要考虑以下几个方面：

1. 开箱即用，低门槛：企业应能快速部署和使用云原生数据库，无需面对过高的技术门槛和成本投入。
2. 跨平台部署：考虑到数据迁移成本较高，企业在选择数据库时会优先考虑数据库的部署位置，如云服务商、自建 IDC、地域等。
3. 兼容性：兼容性对数据库来说至关重要，尤其在当今信息化时代。为了方便企业集成和使用，普惠的数据库需要与各类应用和工具兼容。
4. 稳定性：稳定性是用户选择数据库的重要因素，这不仅关乎数据库的可用性，更关乎数据库的性能表现。

当前，云计算已成为企业数字化转型的重要基础设施，公有云成为云计算的主要形式。然而，随着技术的发展和企业对数据安全、合规的考虑，越来越多的企业选择自建云底座。K8S 技术作为容器编排领域的先锋，为企业自建云底座提供了强大的技术支撑。

面对未来混合云、边缘云等多云环境的普及，企业对云服务的需求将更为多元化。普惠的云原生数据库不应只局限于公有云环境，而应支持多种云服务形态，满足企业在不同场景下的需求。因此，普惠的云原生数据库需要不断拓展应用场景，以更好地满足企业需求。

总的来说，回顾 2023，展望 2024，云原生数据库朝着更高可用，更强处理能力，更智能，更普惠的方向发展。

五、信创云数据库

（一）背景

随着信创产业深入推进和新兴技术的不断迭代，信创云数据库在信息技术创新过程中变得愈发重要。信创云数据库是由信创关系数据库和云计算技术相结合的技术路线发展而来，信创云数据库作为信创行业方案的关键组件，其发展受到了信创行业、技术创新、市场需求、政策环境等多方面因素的影响，目前已呈现出了新的趋势和特点。

本文重在探讨信创云数据库新的发展趋势，分析各信创云数据库提供商在信创行业中的地位以及未来可能面临的挑战和机遇，通过研究信创云数据库的核心技术、技术发展历程以及生态产品，让用户更好理解信创云数据库在信创行业领域的角色和贡献。

（二）发展历程与国内现状

数据库行业近年来经历了巨大的转变，从传统的关系型数据库向更为灵活、高效的云原生分布式数据库技术逐步演进，在这一演进过程中，信创云数据库承担着推动者的角色。

● 数据爆炸和数字化浪潮

随着信息技术的普及和社会的数字化转型，企业和组织生成的数据量呈爆炸式增长，这些数据不仅包括传统的结构化数据，还包括非结构化和半结构化数据。信创云数据库背靠云计算基础设施，能够应对这一庞大且多样化的数据，为用户提供存储、管理和分析的综合解决方案。

● 云原生和微服务架构

云原生和微服务架构在应用开发和部署方面取得了巨大成功。信创云数据库积极响应云原生化的潮流，提供多样的云原生数据库服务，适应了用户对于弹性、可扩展性和灵活性的需求。这使得数据库系统更好地融入容器化、微服务化的应用生态。

● 大数据和实时计算需求

如今大数据和实时计算的兴起，企业对于快速、实时的数据分析和决策支持需求不断增加。信创云数据库不仅支持传统的 OLAP 和 OLTP 工作负载，还能实时数据处理，为用户提供全方位的数据处理能力，以满足业务连续性、降本增效等场景诉求。

- **数据安全和合规性要求**

数据泄露、网络攻击等安全威胁不断加剧，行业对于数据隐私和合规性需求持续提升。信创云数据库在架构和技术层面更加注重数据的安全，可提供丰富的安全功能，以满足用户对数据保密性和合规性的高要求。

- **人工智能和信创云数据库的整合**

当前人工智能和机器学习的广泛应用，对于能够支持智能参数调优、SQL 调优、数据库性能优化、复杂数据分析、模型训练的数据库需求逐渐增长。信创云数据库通过整合人工智能（AIDB+DB4AI）技术，为用户提供更强大的智能化处理能力。

- **多云战略的兴起**

企业越来越倾向于采用多云战略，将工作负载分布在不同的云服务商之间，以降低风险和提高可用性。信创云数据库通过云中立性的设计，使得用户能够在多个云平台之间自由切换，实现高度的灵活性和弹性。

- **数字经济发展**

伴随全球数字经济的迅速发展，企业对于数字化基础设施的需求愈发迫切。信创云数据库通过提供全面的数据库服务，支持企业数字化转型，助力用户在数字经济时代保持竞争力。

- **高效访问和高效响应**

由于信创云数据库对 OLAP、OLTP、实时数据处理的支持，使得性能更加出色，能高效利用存储、计算、网络资源，充分运用 RDMA、NVMe-oF、DPDK、SPDK、CXL、GPU 计算加速等技术，提升或加速信创云数据库的服务能力。

（三）信创领域的特殊性与技术需求

信创云数据库展现出的特殊性和技术需求能满足信创领域的业务特点和技术环境。

首先，信创领域的异构设备与系统的普遍性对信创云数据库提出了高度适应性的要求。由于信创领域涉及多种硬件架构（尤其是非 X86 架构），数据库须具备跨平台、跨设备的兼容性，以确保在各种异构设备和系统上能稳定运行，并提供高效、可靠的数据服务。

其次，信创领域对数据安全和网络隔离的要求更高。由于信创领域中涉及用户的敏感数据和核心业务，数据安全和隐私保护成为首要考虑因素。信创云数据库需要采用先进的加密技术和安全机制，确保用户数据在传输、存储和处理过程中的机密性和完整性。同时，实施网络隔离和访问控制，有效防范网络攻击和数据泄露，满足用户对数据安全的严格需求。

此外，**云中立性和云平台对接**的灵活性也是信创云数据库的重要特性。信创领域的业务通常涉及多个云平台和云服务提供商，用户需根据业务需求自由选择云服务商。因此，信创云数据库需要具备云中立性，能够无缝对接不同的云平台，实现跨云的数据迁移和共享。这种灵活性不仅能够满足用户的业务需求，还能够降低用户的运营成本和风险。

为满足上述特殊性和技术需求，信创云数据库需要具备以下核心技术和能力：

- **跨平台兼容性**

支持多种硬件架构和操作系统，确保在各种异构设备和系统上都能稳定运行。

- **数据加密与安全机制**

采用先进的加密技术和安全机制，保障用户数据的机密性和完整性。

- **网络隔离与访问控制**

实施网络隔离和访问控制，有效防范网络攻击和数据泄露。

- **云中立性与跨云对接**

无缝对接不同的云平台，实现跨云的数据迁移和共享，满足用户的灵活性和扩展性需求。

综上所述，信创云数据库在信创领域中展现出的特殊性和技术需求推动了信创云数据库技术的不断创新和发展。通过具备跨平台兼容性、数据加密与安全机制、网络隔离与访问控制以及云中立性与跨云对接等核心技术，信创云数据库将能够更好地满足用户的业务需求，为信创领域的数字化转型提供强大的数据支持。

（四）信创云数据库技术发展趋势

在信创领域，确保数据库的可用性已经无法满足企业的需求，用户期待数据库能够提供更好的性能、更高的可用性以及更优质的使用体验，因此，信创云数据库在技术的研发上不仅关注可用性，更强调性能优化和用户体验的提升。信创云数据库技术发展趋势主要体现在以下几个方面：

- **云原生**

云原生是信创云数据库技术发展的首要趋势。通过引入云原生理念，实现多点写入、缓存融合、存算分离、算子下推、日志即数据等新兴技术，这些技术能够提供更灵活、可扩展的数据库服务，更好适应云原生和微服务架构的应用需求。云原生将使信创云数据库更加高效、稳定，更好支持企业数字化转型。

● 企业级保障

企业级保障是信创云数据库技术发展的另一关键趋势。以满足企业级需求为目标，信创云数据库在关注数据安全、服务安全、高性能、异地容灾、多活容灾等方面，为企业提供全方位的数据库保障。通过加强数据加密、访问控制、安全审计等措施，信创云数据库将确保企业数据的安全性和机密性。同时，通过优化系统架构、提升处理能力、实现容灾备份等手段，信创云数据库将确保企业的高可用性和稳定性，为企业的持续发展提供坚实的支撑。

● 智能参数调优

当前众多数据库厂商包括信创云数据库，在管理运维方面都存在诸多痛点，例如：数据库设计仍然基于经验方法和人工规则；需要大量的人员参与，问题复杂时可能需要一个或众多 DBA 来调整和维护数据库，时间和成本耗费巨大。

而随着人工智能和机器学习技术的爆发式发展，智能参数调优成为信创云数据库技术发展的新兴趋势。借助人工智能技术，信创云数据库能够自动分析和优化数据库参数，提高数据库的性能和稳定性。通过机器学习算法，信创云数据库能够不断学习和适应业务变化，自动调整参数配置，以满足不同业务场景的需求。智能参数调优将极大降低数据库运维成本，提高业务运行效率。

● 人工智能和机器学习

人工智能和机器学习在信创云数据库技术中发挥着越来越重要的作用。除了智能参数调优外，人工智能和机器学习还可应用于故障预测、性能优化、自动扩容、基于学习模型的优化器及索引推荐及自动视图生成、机器学习索引等方面。通过收集和分析数据库运行数据，人工智能和机器学习能够预测潜在故障并提前进行干预，避免业务中断。

同时，通过对历史数据的分析和学习，人工智能和机器学习能够优化数据库性能，提高查询速度和吞吐量。此外，人工智能和机器学习还可以根据业务需求自动调整数据库资源分配，实现动态扩容和缩容，提高资源利用效率。

数据库优化器依赖于代价和基数估计，但是传统的技术不能提供准确的估计，基于深度学习来估计查询代价和基数，产生基于学习模型的优化器、索引推荐、自动视图生成。传统数据库是由数据库架构师根据经验设计，信创云数据库是基于机器学习，通过

学习的自设计技术，产生学习型索引、学习型数据库设计等，利用人工智能技术来优化数据库，从而使数据库更智能。

信创云数据库技术发展趋势包括：云原生化、企业级保障、智能参数调优以及人工智能和机器学习的应用。这些趋势将推动信创云数据库技术的不断创新和发展，为企业的数字化转型提供更加强大、高效、智能的数据支持。

- **高效率服务能力**

随着硬件的不断迭代更新，也使得信创云数据库有了新的可能。采用新的硬件，加速查询检索过程，实现近数据计算能力。在数据交换过程中，可以达到极低的网络损耗，或者达到网络无损，在网卡和内存中完成通信和计算，降低对 CPU 的依赖和压力。上述众多手段，使得信创云数据库响应更高效。

- **性能至上**

在信创领域，企业对于数据库性能的要求日益提升。为了满足这一需求，信创数据库积极采用了如内存计算、分布式查询优化、智能检索等新技术和优化手段，以提升数据处理的速度和效率。这些技术不仅能够提供信创云数据库的查询性能，还能够优化数据的存储和访问方式，从而为企业提供更快速、更稳定的数据库服务。

- **高可用性和灾备恢复**

在信创云数据库技术的发展中，高可用性和灾备恢复能力也受到了前所未有的关注。通过引入多活容灾、自动故障切换、实时备份等技术，信创云数据库能够在故障发生时迅速切换到备用节点，确保业务的连续性和数据的完整性。同时，通过构建完善的灾备恢复体系，信创云数据库能够在灾难发生时迅速恢复数据和服务，最大限度地减少用户的损失。

（五）信创云数据库典型需求场景和优势

不同行业对数据库多元需求的驱动，信创云数据库也展现出适用于多种场景的特性。尤其在电商、金融等领域。信创数据库以强大的性能、高可用性和可扩展性为特点，为各种业务需求提供卓越的数据服务。

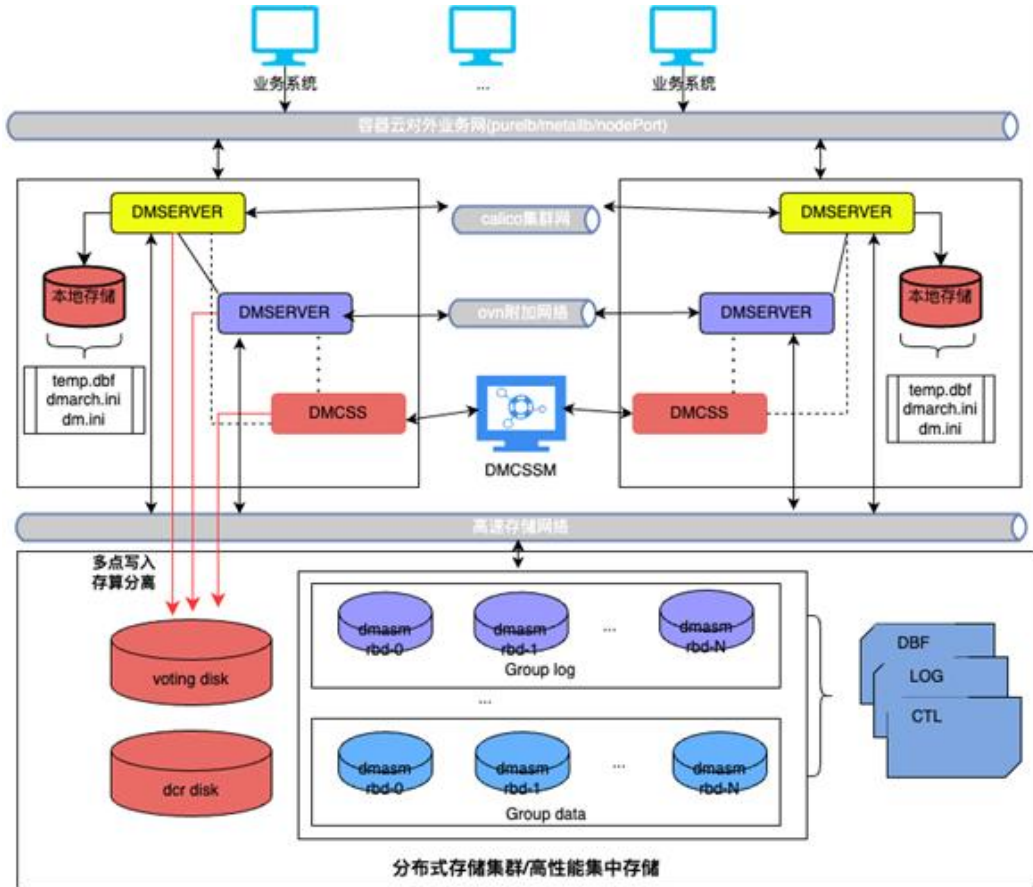
- **信创云数据库多场景需求**

企业在运营中可能涉及不同类型的信创数据库场景，包括守护集群（DWC）、分布式集群（DPC）、读写分离集群（RWC）和共享集群（DSC）等，信创云数据库提供全面支持，能够满足不同场景下的工作负载需求，保证数据库系统在不同应用场景中的灵活运用。

值得一提的是，达梦自主研发的数据共享集群 DSC 和分布式集群 DPC，作为两种高级数据库集群架构，凭借卓越的可用性、灵活的可伸缩性和高效的数据处理能力而备受瞩目。特别是在信创环境下，这两种集群架构展现出了显著的优势，为企业的数据处理和存储需求提供了强大的支持。而在信创环境下，这两种集群所展现出的优势更是尤为显著，为企业带来了前所未有的数据处理和存储体验。

● 数据共享集群 DSC

达梦 DSC 以最高级别的可用性和最灵活的可伸缩性为设计目标。每个 DSC 数据库都是一个由多个独立服务器协同工作的集群数据库，这种架构显著提升了错误恢复能力，并支持逐步扩展。在 DSC 中，客户端通过配置连接服务名轻松连接到集群中的数据库，实现了负载均衡和故障转移的无缝衔接。DSC 集群服务器通过局域网紧密相连，并共享存储资源，确保数据的一致性和持久性。此外，DSC 采用私有专用网络进行内部通信，从而保证了高可用性和可伸缩性。

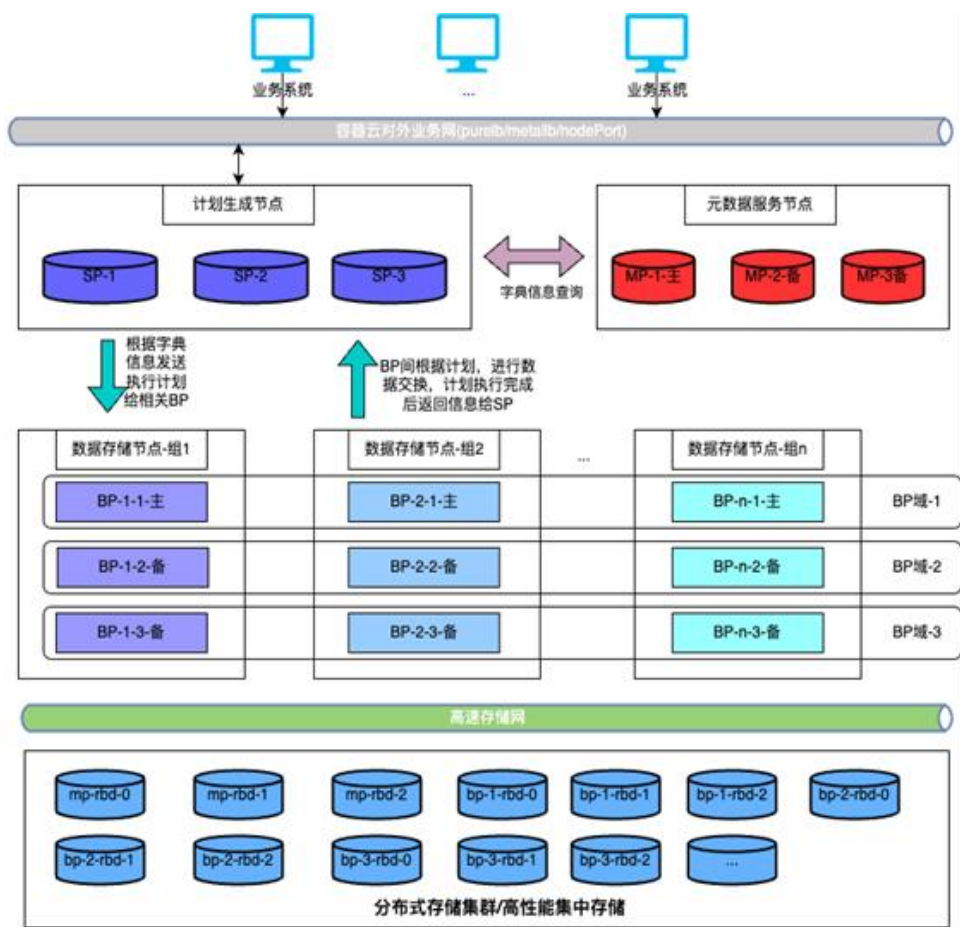


在信创云数据库下，达梦 DSC 集群的优势更加突出，可以在云原生环境下提供多点写入、缓存融合、存算分离等多种差异化特性。DMASM 提供的分布式文件系统不仅支持多节点并发访问和修改数据文件，还消除了对裸设备的依赖，极大地提高了数据处理的灵活性和效率。透明故障切换技术的运用，使得故障转移过程对用户完全透明，确

保了业务的连续性和服务的稳定性。DSC 的水平和垂直扩展能力，使得企业可以根据业务需求灵活调整服务器数量和性能，既满足了当前需求，又为未来发展预留了空间。

● 分布式集群 DPC

达梦 DPC 充分利用多台服务器上的处理能力和存储资源，实现数据的并行处理和存储。在 DPC 中，数据被分割成多个部分并分布在各个节点上，这种分布式存储方式显著提高了数据访问的并行性和吞吐量。计算任务也在各个节点上并行处理，从而充分利用了集群的处理能力，大幅提升了数据处理速度。



在信创环境下，DPC 的优势得到了进一步释放。其强大的容错性保证了即使节点发生故障，系统依然能够稳定运行，数据的完整性和可用性得到了有力保障。通过数据备份、复制和故障转移等技术的综合运用，DPC 实现了高可用性和容错性的完美结合。同时，DPC 还具有出色的可扩展性，可以根据需要动态调整节点数量，轻松应对不断增长的数据处理需求。负载均衡功能则确保了集群中各个节点的资源得到充分利用，避免了资源浪费和性能瓶颈。

达梦 DSC 和 DPC 这两种集群架构在信创环境下展现出了卓越的性能和优势。它们不仅提供了高可用性、灵活的可伸缩性和高效的数据处理能力，还通过先进的技术和机

制确保了业务的连续性和服务的稳定性。这些特点使得 DSC 和 DPC 成为企业应对不断增长的业务需求、提升数据处理和存储效率的理想选择。

● 高安全异构数据库

面对信创领域高安全性的异构数据库需求，信创云数据库提供对 KV、图、文档等多种异构数据库的支持，为用户提供更灵活的数据存储和处理选择。

达梦新云 ForRedis 版本具备兼容开源 Redis 协议的特性，既继承了开源 Redis 的优点，又在自主可控的安全性方面有显著优势，其在并发访问性能上相比原生 Redis 具有明显优势，可满足业务系统对具有高性能、高可靠性、弹性扩展、数据强一致等特性的 Key-Value 系统的需求。因此，在信创云数据库环境下当之无愧的首选。

特别值得一提的是，达梦 GDMBASE 不仅严格遵循图数据库的行业标准，还自主研发了高性能的图存储引擎。这一引擎采纳了业界领先的分布式原生图存储设计理念，确保数据的高效存取与处理。同时，它深度兼容主流的图数据库查询语言，为用户提供了丰富的图应用场景计算能力，从而帮助开发者轻松构建出具有深刻洞察力的应用。

此外，达梦 GDMBASE 还配备了直观易用的系统管理平台，使得应用开发者能够迅速上手并基于图的应用工程进行开发，进而构建出多元化的联系，洞察万事万物。特别是在信创云数据库的环境下，达梦 GDMBASE 展现出了广阔的发展前景，有望为更多企业和应用带来更高效、更智能的数据处理与分析能力。

● 运维巡检智能化 AI

随着数据库规模的不断扩大和复杂性的增加，传统的运维方式已经难以满足高效、精准的管理需求。为此，信创云数据库积极引入智能化 AI 技术，对运维巡检进行革新。通过智能化的性能监测、故障诊断和自动优化调整，信创云数据库能够实时感知系统的运行状态，预测潜在风险，并自动进行干预和调整，确保数据库系统始终运行在最佳状态。这不仅极大地提高了运维效率，减少了人工干预和误操作的可能性，还为企业的业务连续性提供了坚实保障。

● 集约化监控巡检

为了实现对数据库系统的全面、高效管理，信创云数据库采用集约化监控巡检策略。通过整合各类监测工具和巡检手段，信创云数据库能够实现对数据库系统的全方位、多层次监控，确保系统的各个方面和细节都得到了及时地关注和处理。这种集约化的管理方式不仅提高了监控效率，还降低了运维成本，为企业带来了更大的价值。

● 成熟产品上云和多形态支持

随着企业业务的快速发展和数字化转型的深入推进，企业对数据库的需求也日益多样化。为了满足不同规模、不同形态的业务需求，信创云数据库提供了成熟产品上云和多形态支持的服务。无论是大型企业的大型案例，还是中小企业的中小型案例，无论是容器、物理机还是云管服务等不同部署场景，信创云数据库都能够提供全面、高效的支持。这种灵活、多样化的服务方式不仅满足了企业的实际需求，也为企业的发展提供了强大的技术支撑。

信创云数据库通过引入智能化 AI 技术、采用集约化监控巡检策略以及提供成熟产品“上云”和多形态支持的服务方式，不断提升自身的技术实力和服务能力，为企业的数字化转型提供了强有力的支持。

（六）信创云数据库未来展望

未来，信创云数据库将聚焦于技术创新和升级，以满足不断变化的业务需求和技术趋势。

● 技术创新和升级

技术创新是信创云数据库持续发展的核心驱动力。未来，信创云数据库将继续投入研发力量，推动技术创新和升级。通过引入新的算法、优化数据库架构、提升性能等方式，信创云数据库将不断提升自身的技术实力，为用户提供更加高效、稳定、安全的数据服务。

● 全球拓展和服务地域的扩张

随着全球化的加速推进，信创云数据库将积极拓展全球市场，扩大服务地域的覆盖范围。通过与全球合作伙伴的紧密合作，将提供更多的本地化服务，满足不同地区、不同文化背景下用户的需求。这将有助于提升信创云数据库的品牌影响力，同时也为企业带来了更大的发展机遇。

● 数据安全和合规性

在数据安全日益受到重视的背景下，信创云数据库将进一步加强数据安全和合规性方面的投入。通过加强数据加密、访问控制、安全审计等措施，信创云数据库将确保用户数据的安全性和机密性。同时，信创云数据库还将积极遵守各国法律法规，确保业务的合规性，为用户提供更加可靠的数据服务。

● 用户体验的提升

用户体验是信创云数据库持续优化的重要方向。未来，信创云数据库将更加注重用户体验的提升，通过优化界面设计、简化操作流程、提供智能化支持等方式，降低用户

使用门槛，提升用户使用满意度。这将有助于增强用户粘性，推动信创云数据库的市场份额不断扩大。

- **生态系统的建设**

生态系统的建设是信创云数据库未来发展的关键一环。通过与各类合作伙伴的紧密合作，信创云数据库将构建一个完善的生态系统，为用户提供更加全面、丰富的数据服务。这将有助于提升信创云数据库的市场竞争力，同时也为合作伙伴带来了更多的商业机会。

综上所述，信创云数据库的未来展望充满了机遇和挑战。通过技术创新和升级、全球拓展和服务地域的扩张、数据安全和合规性、用户体验的提升以及生态系统的建设等多方面的努力，信创云数据库将不断提升自身实力，为用户提供更加高效、稳定、安全的数据服务。

2、HTAP

一、背景

随着企业数据规模的不断增长和对实时性的迫切需求，数据库技术也在不断演进。HTAP 数据库（Hybrid Transactional/Analytical Processing）应运而生，作为一种能够同时支持事务和分析处理的数据库技术，为企业数据管理带来了全新的可能性。

二、HTAP 数据库的优点

- **节省时间和成本**

HTAP 数据库的一项显著优势是，无需将数据从操作型数据库导入到决策类系统，从而节省了时间和成本。

- **实时可见性**

操作事务实时地对分析业务可见，提高了数据新鲜度和分析效率，使企业能够更迅速地做出决策。

- **减少数据冗余**

HTAP 数据库减少了对副本的要求，降低了数据冗余和不一致的风险，提高了数据的一致性。

- **弹性扩容**

HTAP 数据库支持弹性扩容，可以按需扩展吞吐或存储，轻松应对高并发、海量数据场景。

三、HTAP 数据库的发展历程

- **第一阶段（2010-2014）**

在初始阶段，HTAP 数据库主要采用主列存的方式，代表性产品有 SAP HANA、HyPer、DB2 和 BLU 等。

- **第二阶段（2014-2020）**

HTAP 数据库分布式环境下，扩展了主行存的技术，在行存上加入了列存，代表性产品有 SQL Server、Oracle 和 L-store 等。

- **第三阶段（2020-至今）**

HTAP 数据库开启了存算分离，云边一体的分布式架构实现，以满足容器化趋势下，高并发的请求，代表性产品有 SingleStore、Tidb、MatrixOne 等。

四、HTAP 数据库的现状

（一）国际发展

截至 2023 年，HTAP 数据库在国际市场上经历了显著的发展。全球范围内，企业对实时数据分析的需求日益增长，推动了 HTAP 技术的广泛应用。多个国家和地区的企业都在寻求这一新型数据库技术以更好地满足他们的业务需求。

主流的国际数据库供应商积极响应 HTAP 数据库的需求，推出了相应的解决方案和产品。例如，美国的 Oracle 和 Microsoft、德国的 SAP 等公司在 HTAP 数据库领域取得了显著的成就。这些公司通过不断创新和改进，提供了更高效、更可靠的 HTAP 数据库产品，满足了各行业客户对实时性和性能的迫切需求。

HTAP 数据库在国际上得到了广泛的行业应用。金融领域通过实时交易数据和风险分析提高了决策速度；制造业通过实时监控生产过程优化生产效率；零售业通过实时分析购物行为提升了用户体验。这些应用案例表明 HTAP 数据库在各个行业都取得了显著的成就。

随着 HTAP 数据库技术的逐步成熟，国际上形成了一定的竞争格局。各大数据库供应商之间展开激烈的竞争，通过不断推陈出新、提升性能、优化用户体验来争夺市场份额。同时，一些新兴公司也在 HTAP 数据库领域崭露头角，为国际市场的 HTAP 数据库创新注入新的活力。

在国际发展过程中，HTAP 数据库领域的技术合作和标准化变得日益重要。不同国家和组织之间的合作促进了技术的跨界应用和经验的分享。此外，制定统一的标准有助于确保不同供应商的 HTAP 数据库产品之间的互操作性，提升整个行业的发展水平。

未来，随着全球数字化进程的深入推进，HTAP 数据库在国际市场上将继续发挥关键作用。不断提升的硬件技术、智能化的数据管理以及更加成熟的 HTAP 数据库产品都将推动这一技术在国际市场上迎来更大的发展空间。国际合作、创新和标准制定将是推动 HTAP 数据库进一步发展的关键因素。

通过国际市场的介绍，我们可以看到 HTAP 数据库在全球范围内受到了广泛关注和应用，为不同国家和行业带来了显著的价值。

（二）国内市场

中国国内市场在数字化转型和信息化建设的浪潮下，对数据库技术的需求日益增长。HTAP 数据库技术作为同时支持事务和分析处理的创新性解决方案，在国内市场获得了广泛关注和应用。在这一市场环境下，数据库国产化和自主创新成为行业发展的关键驱动力。

数据库信创背后的背景主要包括国家信息安全的考量和降低对外部技术依赖的需求。中国政府提出“自主可控”理念，强调在关键技术领域实现自主创新和掌控，特别是在数据库领域，以确保国家信息安全和数据主权自主可控。国产自主数据库的背景体现了企业对自主创新的渴望。随着数字化时代的到来，企业对于数据库技术的依赖程度增加，因此，具备自主研发能力的数据库系统能够更好地适应不同行业的需求，提供更贴近本土企业实际情况的解决方案。

为推动数据库国产化，提出了一系列要求：

● 技术自主创新

推动数据库国产化需要依托技术自主创新，涉及数据库核心技术的研发和创新，包括存储引擎、查询优化、分布式架构等方面的创新。

● 安全可控

国产数据库必须具备高度的信息安全性，能够保障用户数据在存储、传输和处理过程中的安全，以满足政府、企业对信息安全的严格要求。

● 本土化适配

考虑中国的特殊文化、法规和业务环境，国产数据库需要具备更好的本土化适配能力，以更好地服务国内企业的需求。

● 兼容性与互操作性

国产数据库需要具备与国际标准的兼容性和互操作性，以便在全球范围内广泛应用，实现与国际数据库系统的对接与协同。

国内市场，新老数据库厂商竞争激烈，老厂商代表 - 达梦，人大金仓

老一代数据库厂商中，达梦数据库以其在国内市场的长期积累和丰富经验成为代表。达梦数据库通过持续优化产品性能和提供全方位的数据库解决方案，保持在传统领域的竞争力。中生代代表 - 阿里、腾讯互联网背景下成长的中生代数据库厂商中，阿里云和腾讯云通过改造开源产品，实现云端化和服务化，以云为代表，迅速崛起。这两家公司在数据库领域的开源产品和云服务上进行了深入的布局，形成了与传统厂商的竞争态势。新锐代表 - 矩阵起源、TiDB 新锐数据库厂商中，矩阵起源和 TiDB 以完全自主、云边一体的技术路线为代表。矩阵起源以其独特的分布式存储技术和云边一体的解决方案，满足了不同场景的需求。而矩阵起源则凭借其存算分离分布式数据库的开创性设计，吸引了众多用户。

不同数据库厂商的技术路线存在显著差异。传统厂商如达梦更注重传统关系型数据库的优化和性能提升，中生代厂商如阿里、腾讯通过改造开源产品实现云端化。而新锐厂商矩阵起源和 TiDB 更强调分布式、云原生和云边一体的创新技术路线。

未来，随着数字经济的不断发展，数据库技术将在国内市场扮演更加重要的角色。数据库国产化和自主创新将成为中国数据库行业的主题，新老厂商将在技术、服务和解决方案上展开激烈的竞争。随着新兴技术的涌现，数据库领域将持续迎来创新和发展的机遇。

五、HTAP 数据库面临的挑战

● 数据组织

如何在行存和列存之间实现高效的数据转换和存储,是 HTAP 数据库面临的一个重要挑战。

● 数据同步

在行存和列存之间保持数据的一致性和及时性是 HTAP 数据库需要克服的问题之一。

● 查询优化

根据不同的工作负载选择合适的存储引擎和执行计划,是提高 HTAP 数据库性能的一项关键挑战。

● 资源调度

在有限的资源下平衡事务和分析的性能和隔离是 HTAP 数据库需要解决的复杂问题。

六、HTAP 数据库的未来趋势与创新

● 利用新兴硬件技术

HTAP 数据库未来趋势之一是利用新兴硬件技术，如非易失性内存 (NVM)、GPU、FPGA 等，以提高性能和可扩展性。

● 机器学习和人工智能的整合

HTAP 数据库将整合机器学习和人工智能技术，实现自适应查询处理、自动索引选择、自动参数调优等功能，提高智能化和自动化水平。

● 云计算和边缘计算技术的融合

HTAP 数据库未来将充分利用云计算和边缘计算技术，包括云原生架构、多租户支持、异地多活等，以提高可用性和可移植性。

● 流式计算和时序数据库技术的整合

HTAP 数据库将整合流式计算和时序数据库技术，实现流批一体、实时分析、时序查询等功能，提高实时性和时效性。

总结：

综合而言，HTAP 数据库在国内外都取得了显著的进展，不断迭代以适应企业不断增长的需求。面对挑战，HTAP 数据库不仅不断优化现有功能，还积极探索新技术的应用，为企业提供更强大、智能、高效的数据管理解决方案。随着技术的不断进步，HTAP 数据库将持续发挥关键作用，推动企业在数字化时代取得更大的成功。

3、共享集群

一、背景介绍

共享集群数据库管理系统是一种单库多实例的多活数据库管理系统，用户连接任意实例都可以访问同一个数据库，多个数据库实例可以并发读写同一份数据，具备高可用、高扩展、高性能等特性。

数据库集群软件根据侧重的方向和试图解决的问题划分为三大类：负载均衡集群（Load balance cluster, LBC）侧重于数据库的横向扩展，提升数据库的性能；高可用性集群（High availability cluster, HAC）侧重保证数据库应用持续不断；高安全性集群（High Security cluster, HSC）侧重于容灾。

共享集群被称为数据库领域的“塔尖”技术，其开发难度高。该核心技术之前一直被国外垄断，主要掌握在 Oracle、IBM 等国际巨头手上，国产数据库厂商拥有该技术的凤毛麟角。

最早于 1990 年 IBM 就提出 Systems Complex（也即 SysPlex）概念，1994 年提出 Parallel Sysplex 概念，并行系统耦合体是大型机最具代表性的集群技术。DB2 在 2009 年推出了基于小型机的 pureScale 集群，但是没有取得像 Oracle RAC 那样的成功。

Oracle 从 90 年代开始尝试开发 Oracle 并行服务器（OPS, Oracle Parallel Server, 为 RAC 前身），Oracle 9i 正式推出 RAC 形态，至今经历 20 多年的发展，Oracle RAC 在金融、电信、能源、交通等行业领域深度应用，尤其在关系国计民生的核心业务系统中，大量采用小型机 + Oracle RAC 组合，依靠系统性能和高可靠方面的优势取得巨大市场成功。

目前国产数据库大多基于 MySQL、PostgreSQL 等开源数据库二次改写，受限于存储结构，其架构向多活共享集群方向演进时非常困难。同时由于共享集群数据库的技术难度，很多厂商产品技术上很难达到对等替换 Oracle RAC 的能力，产品能力还在摸索以及不断追赶阶段。所以到现在为止，尤其是金融等核心系统还是 Oracle RAC 的天下。

在当前的国际背景下，针对该领域的国产化趋势为国产数据库厂商带来千载难逢的机遇，能否把握机会成功取代 Oracle RAC 生态，这也是现阶段国产数据库厂商要回答的问题。

二、传统架构不足及痛点

集中式数据库高可用方案常采用一主多备的方式，通常有同步复制和异步复制两种，如果选择异步复制，故障切换时数据丢失量 RPO 不为 0，如果选择同步复制往往可用性大大降低，略微的网络抖动导致延迟会大大增加，此外无论选择哪种复制方式，单点故障后都会产生额外的运维工作。

而分布式数据库系统往往带来应用架构以及运维的复杂化，一是传统业务搬迁成本巨大，往往需要有较强开发能力的团队才能驾驭；二是分布式常面临三节点起步价格较高、性能劣化等问题。

无论是传统的集中式主备数据库还是分布式数据库，不能实现真正应用透明的多写，副本数据最多提供只读服务，无法充分利用硬件资源的潜力。分布式以及传统主备数据库无法对 Oracle RAC 等共享集群产品对等替换，有些客户甚至用数十台服务器的分布式数据库去替换一套 Oracle RAC 双机，硬件成本和运维成本都大大增加。

三、共享集群数据库优势

为了解决传统数据库的不足，以 Oracle RAC 为代表的共享集群数据库产品通过共享存储数据的设计解决了集中式数据上的主备同步实时性、单节点故障、数据丢失、系统资源利用不足等关键核心问题，并且提供了高可靠、高可用以及易扩展的能力。具体优势表现在：

- 真正对应用透明的多写集群，集群所有节点均可读写全部数据，数据库服务器需要水平扩容时，可以不修改软件，大大降低应用软件架构设计的难度。
- 成熟的高可用架构，经过二十多年的广泛应用，RAC 的高可用价值为广大运维人员所认可，单点故障发生时可自动接管，客户端透明切换，事后复原操作也非常容易，无须考虑重建副本等复杂的运维动作。
- 优良的可靠性表现，故障切换 RPO 为 0，对于无法容忍数据丢失的应用尤为关键，例如金融交易、电信计费场景，RTO 分钟级甚至秒级，保证核心业务的影响最小化。
- 性能优势明显，作为集中式数据库往往更能充分发挥硬件潜力，电信金融等很多核心业务场景往往都是一套或若干套 RAC 双机即可支撑，替换一套 RAC 往往需要数倍甚至数十倍的服务器才能满足性能要求，另外共享集群数据库采用单库多实例的架构，只需要一份持久化数据，存放在共享磁阵上保证存储高可用，避免了存储空间要求的放大，也减少了系统整体 IO 吞吐。
- 良好的可扩展能力，共享集群数据库也具备线性扩展能力，两节点集群即可满足大部分应用要求，交易型场景一般可线性扩展到十节点集群，分析型可线性扩展到上百节点集群。

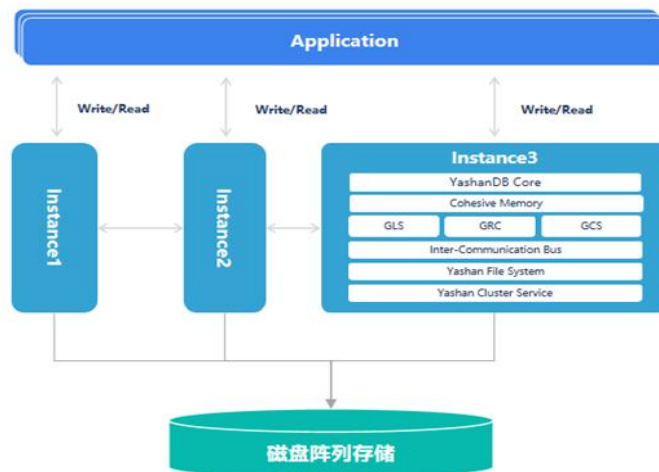
四、国产厂商共享集群案例

由于企业级核心系统对业务连续性要求颇高，希望故障恢复所需时间短、数据丢失量接近为零，因此能够支撑金融级高可用、性能优异的集群数据库备受青睐。

（一）YashanDB for Cluster

目前国产厂商中，2023 年 YashanDB 推出完全自主研发设计、对标 Oracle RAC 的 YashanDB for Cluster。采用单数据库多实例架构，所有计算节点提供对等的多活计算能力，节点之间以强一致性方式实现并发读写，可为用户提供透明多写、高可用、高性能数据库能力，为关键核心业务实现国外数据库“1:1”平替。

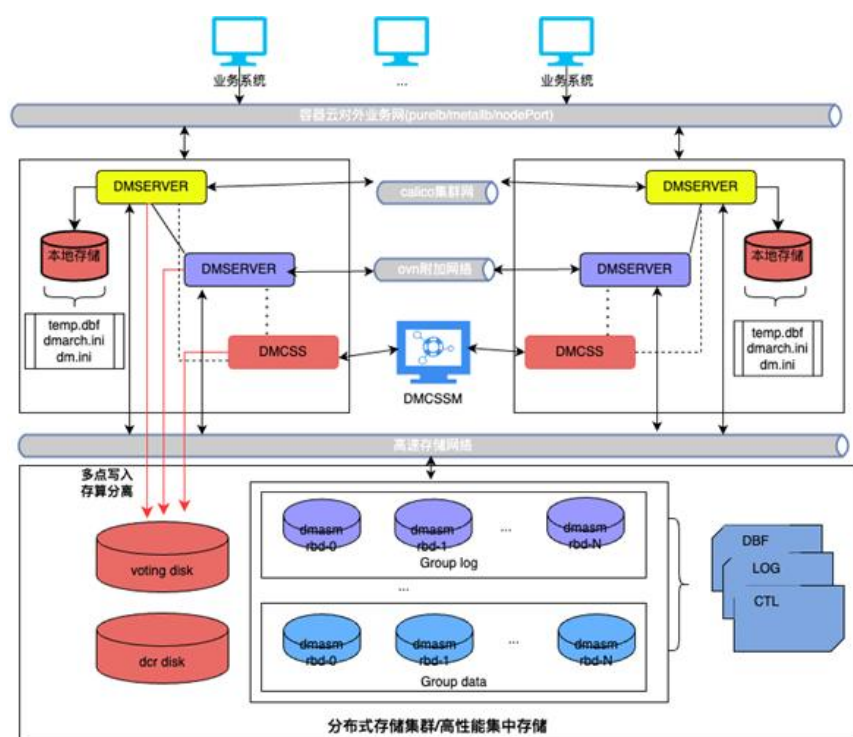
基于磁盘阵列的共享集群数据库



YashanDB for Cluster 架构图

（二）达梦数据共享集群 DSC

达梦 DSC 以最高级别的可用性和最灵活的可伸缩性为设计目标。每个 DSC 数据库都是一个由多个独立服务器协同工作的集群数据库，这种架构显著提升了错误恢复能力，并支持逐步扩展。在 DSC 中，客户端通过配置连接服务名轻松连接到集群中的数据库，实现了负载均衡和故障转移的无缝衔接。DSC 集群服务器通过局域网紧密相连，并共享存储资源，确保数据的一致性和持久性。此外，DSC 采用私有专用网络进行内部通信，从而保证了高可用性和可伸缩性。



4、超融合数据库

一、行业数据应用面临的痛点

看到这个标题后，不少人立即冒出一串疑问，什么是超融合数据库，为什么需要超融合数据库，超融合数据库如何实现....面对这些问题我们不着急回答，先探讨另外一个问题，当下的各行业的数据应用都面临了哪些痛点。

在说痛点之前我们先说说痛点的背景，当下，已经进入物联网+数智化时代，其最大特点是实现了物与物、人与物之间的全面连接和智能化，使各种设备、传感器和人类能够实时共享数据、互动和控制。在此背景下，数据库产生了三大压力，第一个压力来自**数据规模**，当数据从以人为主体的迈向以人+物+时空的联合主体后，其数据规模可能会呈指数级增加；第二个压力来自**用户规模**，当用户从人走向人+物+事件后，其用户规模也必然增大，即并发增大；第三个压力来自**用户需求**，当需求更加强调物理世界与数字世界之间实时、智能的深度融合时，必然会促使用户需求更加复杂多样化，包括处理数据的类型更广泛，数据分析的维度更多元等。

于是数据应用痛点出现了，当数据和用户规模压力上来了，数据的写入、更新和查询就开始遇到挑战了，而需求压力则导致更新查询复杂度不断提升，让其雪上加霜，因此**性能需求难以满足**成为第一个显著痛点。当然，需求的实现也不是一件容易的事，包括预测性维护在内的预见性需求、追溯在内的复杂关联关系需求、车联网在内的时序需

求...而这些数据的类型多样化还进一步提升了难度，因此**功能需求难以实现成为行业数据应用的第二个显著痛点**。

怎么办？对此，行业客户和数据库厂家各显神通，也找到了应对的办法。

二、企业应对数据痛点的方案

先说行业客户，他们方案为六字诀：**限制、升级、拆分**。思路为，既然痛点源自三大压力，咱就分而破之。

先看限制，数据规模太大怎么办？那就给数据来一个瘦身，比如考虑历史数据归档转移等思路。并发数量太大怎么办？那就限制应用的访问，首先是限制**访问资源**，包括 CPU、内存的使用份额。接着是限制**访问时间**，比如高峰期间只允许某些特定的应用访问数据库。需求复杂多样怎么办？那就限制**业务需求**，比如在不影响业务正常开展的前提下，将部分需要实时统计的需求调整成非实时的 T+1。此外还有限制 SQL 复杂度，譬如只允许最多 5 张表 JOIN 等。

再看升级，即硬件与软件的升级。**硬件升级**如使用高性能服务器、扩大内存、使用 SSD 存储，提升网络带宽等。**软件升级**诸如数据库、操作系统、中间件的版本升级，以及 SQL 优化、将数据处理逻辑转移到应用程序端去实现等。

最后看看拆分，即根据**业务模块拆分**数据库。比如把某个业务拆分成若干个不同的子业务，分别用独立的数据库来承接，这种拆分将数据从物理上独立出来，一方面解决数据扩展问题，另一方面能让资源的分配变得更加有效，从而缩小故障影响面，以保障核心关键业务。

那数据库厂家如何应对呢？两字诀：拆分。这与前面提到的针对业务模块拆分有所不同，而是依据**数据模型拆分**数据库，以应对需求难以实现的场景。如：为应对序场景对写入和实时分析的高要求，研发了时序数据库；为应对特定场景对数据处理响应时间的高要求，研发了内存数据库；为应对特定的复杂关联关系分析场景，研发了图数据库...这样，在数据库厂家的支撑下，行业客户将依据数据模型拆分的方案纳入其中，从而让数据处理变得更为精细。

遗憾的是，上述方案在解决旧痛点的同时，也带来了新的痛点。比如升级和使用专用数据库带来了**成本问题**，比如拆分后产生大量数据交互带来了**稳定性问题**。而且无论限制、升级还是拆分，都是有尽头的，不能无限操作，收益有限。不过，而更闹心的是，旧痛点或将卷土重来。比如限制会影响用户的效率，其实变相地降低了系统的性能。比如拆分导致业务处理的任务链条加长，导致实际性能反而更低。比如专用数据往往通过

牺牲 ACID、SQL 等代价来换取特定能力，当场景需要这些被牺牲的能力时，往往需要付出更大的代价。

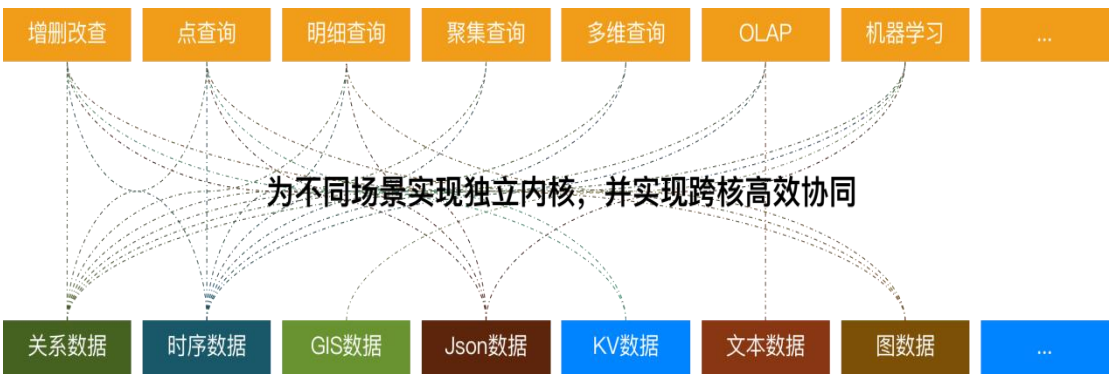
怎么办？更好的方式是为企业提供数据的一站式解决方案来真正解决问题。

三、超融合数据库理念与实现

（一）能力融合

超融合架构如何解决传统数据库的“信息孤岛”的问题，实现“一库多用”？YMatrix 是通过其先进的微内核架构来实现。

这是一种能跨核高效协同，为不同场景提供专门定制的内核，并确保它们之间互不干扰。这种设计使得每个内核都可以最大程度地发挥其优势，提高整体性能。而跨核高效协同则意味着，在需要时，这些内核能够快速协同处理数据，从而进一步提高系统的整体性能。通过极限独立内核和跨核高效协同，超融合数据库成功地整合了多模数据、多模操作、高性能和易用性。



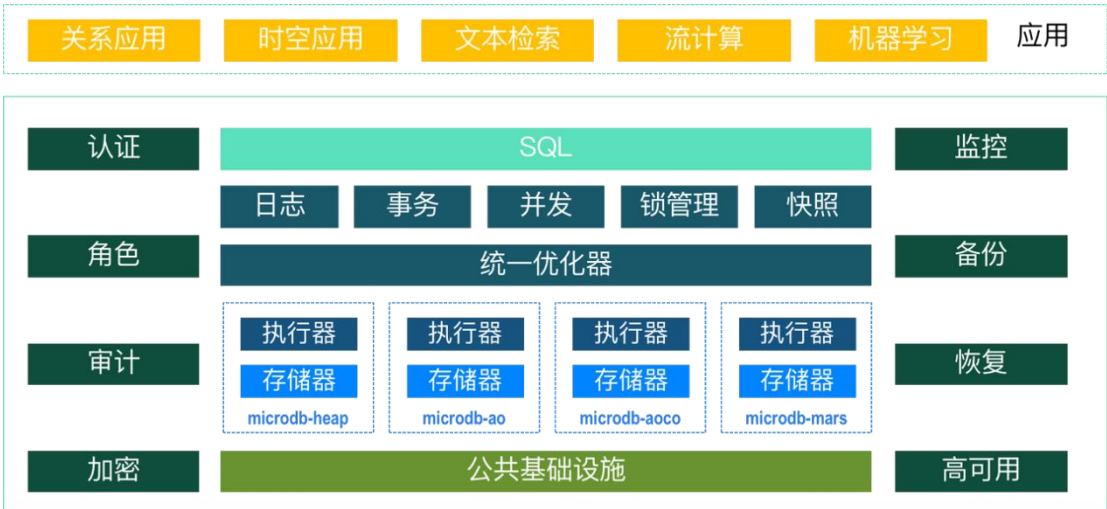
微内核落到具体的数据库设计中，最关键的是在**优化器、执行器、存储器**等这些数据库的内核模块上，这是极具挑战的。我们先从存储器这个层面说起，对于超融合场景，我们需要将存储器设计成可插拔形式，这样不同的场景对应不同的存储器。比如应对 OLTP 场景需要行式存储，我们就设计行式存储引擎。比如应对 OLAP 场景需要列式存储，我们就设计列式存储引擎。比如时序场景需要满足高性能写入和低延迟分析，我们就设计专门针对时序场景优化的 MARS 存储引擎.....接着说说执行器，这也是非常灵活的可插拔设计模式，和存储引擎可以形成各种绝妙的组合来解决各种棘手的问题。比如**向量执行器和列式存储结合，就可以实现数据库的向量化技术**，在某些场景下，性能甚至能有百倍的提升。最后说说优化器，微内核选用的存储引擎通常是固定的，执行器则根据优化器代价评估的结果而定。例如，在处理大量复杂查询的 OLAP 场景中，我们可能会选择一个专门针对列式存储优化的执行器，以便更有效地进行数据分析和聚合操

作。反之，在 OLTP 场景中，由于需要频繁进行行级操作，如更新和插入，执行器可能会被优化以更适应行式存储，从而提高事务处理的效率。

微内核则是超融合实现的最关键、最核心之处，是优化器的统一管理下，通过存储引擎和执行引擎之间的巧妙配合，不同的微内核为不同的场景而优化。如：通过 HEAP 存储引擎 + 火山执行引擎实现 OLTP 微内核，以应对 OLTP 场景需求；通过 MARS3 存储引擎+火山执行引擎实现 OLAP 微内核，以应对 OLAP 场景需求；通过图存储引擎+图执行引擎的组合实现图微内核，以应对复杂关系的图场景需求；通过 MARS3 存储引擎 + 向量化执行引擎实现时序微内核，以应对海量写入和实时分析的时序场景需求等等。此外，场景的实现不一定只是某种存储引擎和某种执行引擎之间的 1+1 组合模式，还可以是多种存储引擎和多种执行引擎的 N+N 的组合模式。比如在 OLAP 微内核的基础上增加向量化执行引擎,变成 AOCO 存储引擎+(火山执行引擎+向量化执行引擎)的 1+2 组合模式，就能应对更极致的 OLAP 场景；比如将 OLAP 与 OLAP 微内核的实现进行巧妙结合,变成(HEAP 存储引擎+AOCO 存储引擎)+(火山执行引擎+向量化执行引擎)的 2+2 组合模式，就能解决 HTAP 场景需求.....微内核的选择完全无需人工来干预，真正做到把复杂留给数据库本身，把简单留给用户。

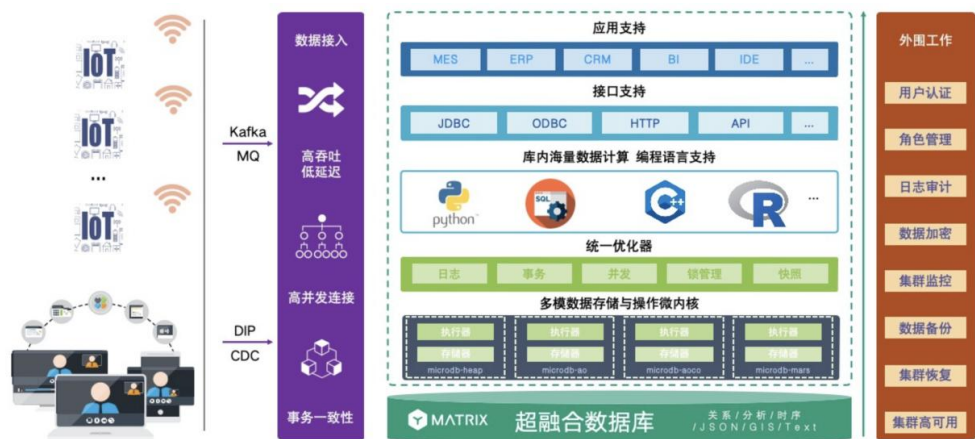
(二) 接口统一

当然了，微内核架构离不开**统一 SQL 层面的配合**，以时序场景为例，比如在查询滑动窗口、时间间隔、最新值等需求时，如果通过传统的 SQL 语句来实现，不仅书写难度大，性能也难以保障，这就要求超融合数据库需要从内置高性能的自定义函数，从而更方便更高效地实现这些复杂需求，类似的情况不仅在时序场景，还有很多其他场景也类似，如流处理场景、库内机器学习、图场景等也都是如此，在同一优化器下执行统一的 SQL 语言，完成高效交互。



（三）开放架构

除了能力的融合和接口的统一，提供了丰富的 API 和工具集，促进了与其他系统和应用的集成，打造开放的架构也是必不可少的。比如必须支持标准化接口协议，使得不同数据处理场景下的用户交互变得更为便捷；支持与 Kafka 等消息队列系统的集成，使数据平台能够更便捷地处理实时数据流和大数据场景；支持数据联邦，允许跨数据库、跨平台的数据共享和交互，增强了企业间的数据合作能力；整合库内机器学习能力，简化了数据处理流程，降低了机器学习项目的门槛，还能极大地提升数据分析处理的性能。



四、超融合数据库的经典案例

超融合数据库的好处自是不言而喻，接下来我们来看一个制造业巨头的企业大数据平台案例，该平台承载了销售预测、设备预测性维护、质量大数据预测、质量/成本追溯、生产运营管理、数字孪生等模块，业务场景非常复杂。

（一）取得成效

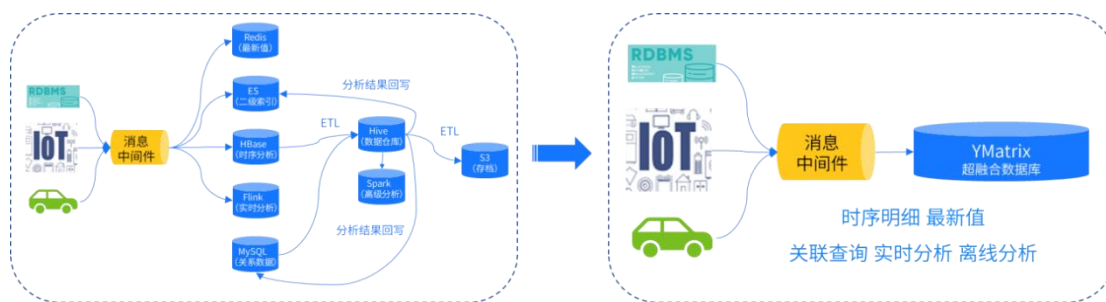
当前该平台存在如下痛点，1. 稳定性：出现宕机的情况，这是企业用户最无法接受的；2. 功能：无法实现实时预测，无法应对乱序、异频上传的时序数据等，导致项目有所延误；3. 性能：写入速度不足，查询性能无法满足业务需求，尤其是追溯场景无法满足需求；4. 成本：由于技术栈繁多，导致采购、人力和时间成本大幅增加。

经过超融合数据库 YMatrix 的改造后，取得了意想不到的效果。从稳定性来看：整体 CPU 使用率降低 25% 左右；宕机问题不再发生；从功能来看：实时预测得以实现，其库内机器学习性能达到 Spark 模式的 10 倍以上；时序各种异常数据均得到有效地处理，且给用户提供了更为丰富灵活的查询能力；从性能角度来看：整体写入和查询性能均提升 5 倍以上，追溯场景的整体处理能力提升 5 倍以上；从成本来看：整体节省费用超 3000 万元，人力资源节省 50%，开发周期缩短一半以上。

（二）成效剖析

这是如何做到的呢？

先来看看客户最关注的稳定性问题，请看超融合改造的部分架构图，其中右图为改造后的技术栈，架构化繁为简，清爽简洁。原先的宕机定位为系统间的频繁交换相关，现在交互大幅减少了，数据库节点莫名宕机的问题也就自然被修复了，这为企业长吁了一口气。



功能何以得到有效支撑？说说实时预测的实现。原先机器学习性能大幅提升也是减少交互的典型，超融合数据库将机器学习的能力内嵌到数据库里，我们称之为库内机器学习。库内机器学习让计算贴近数据，无需搬动数据，对比原先 Hadoop+Spark 的库外机器学习模式则数据需要出库处理，在性能上有明显的优势。加之库内机器学习还能有效地应用到分布式数据库的并行能力，进一步提升计算效率，最终库内机器学习比 Spark 模式性能提升 10 倍以上，满足了实时预测的需求。

那性能何以大幅提升，我们来从场景优化中寻找答案。该企业为快速定位产品的复杂关联关系，选择了专用图数据库，以便更好地寻找有质量问题的环节与参数，原本企业希望这样能有更好的性能，实际却不尽如人意。因为这是一个复杂的组合场景，不仅是要追溯问题出在哪个环节，还要立即对有问题环节进行围堵，以避免其进入下一环节，同时还要做更大范围的影响评估，包括关联影响和供应链追查等。如此一来数据交互必不可少，专用数据库的不足被放大了。此外由于该业务场景的追溯层级基本固定的，关系模型数据库的性能并不比图数据库低多少。经过评估，放弃图数据库组合使用通过关系模型 SQL 能力成为更好的选择，最终整体性能提升 5 倍以上。

关于成本这块我们仅说开发周期为啥能减少。在超融合架构下，预测分析类需求使用库内机器学习模式，缩减了算法难度，加之由于库内可以用到分布式并行能力，使库内机器学习的效率有了极大提升，也大幅提升了数据测试、验证的时间。此外大量时序类数据由于融合了预聚集、持续聚集及强大的窗口函数等特性，也大幅提升了后台开发的工作。

五、超融合数据库思考与展望

至此，说完了超融合数据库的理念、实现与经典案例，接下来我们抛出一个问题来探讨：未来，超融合数据库会是一种主流的存在吗？

这个问题隐藏着另外一层含义，就是超融合数据库当下并非主流，是的，当下确实并非主流。有人会问，这又是为何，超融合数据库不是很好吗，为啥当下无法成为主流？问得好，我们先来回答这个问题。

影响历史进程的因素是综合性的，并非具备条件就一定会完成替代。乱世出英雄，话说专用数据库在互联网时代得以快速发展后，很快就承载了大量应用和各种需求，同时培育了一大批技术专家....此时要想替换就不太容易了，毕竟圈子形成后替换的成本是很高的。那随后功能不满足需求该咋办？那就考虑增加新的技术栈来组合实现，或者转移到应用程序端去实现。有趣的是，反观关系型数据库，虽然早已具备了融合各专用数据库的良好基础，却并未对收编专用数据库展示出过多热情，为啥如此，主要是因为彼时的关系型数据库已经承载了大量应用，而统一整合意味着要对现有所有用户的数据库版本进行升级改造，必然存在风险。而保持关系型数据库现有架构不变，通过对接专用数据库来弥补其缺失功能的方式，无疑成为最为安全的选择。以关系型数据库的绝对王者 Oracle 为例，就接连推出了支持时序、图、GIS、内存、分析等一众独立专用数据库，让人眼花缭乱。以 Oracle 的实力绝非不能推出真正意义上的超融合数据库，而是船已大，不敢推，这已超出技术范畴了。

如此说来，似乎超融合数据库一直不会成为主流的存在了，并非如此。在社会发展中，经济规律是不会被轻易动摇的，任何产品最终都会朝着价值利益最大化的方向发展，否则就会被淘汰。超融合数据库无论从哪个角度上，收益都是最高的，所以将成为数据库发展的一种必然趋势。回想当年柯达已经掌握了全球最领先的数码相机的技术，却因为众所周知的原因固守着旧阵地，由于数码相机的对人们的价值收益远超胶卷相机，最终胶卷相机就被历史的车轮滚滚淹没了，当然数码相机最后也融入了智能手机中，向对人们价值利益更大的方向迈进了。

未来，超融合数据库能做的事将越来越多，超融合数据库将逐步走向主流，这会是随着时间推进自然而然发生的事。超融合数据库好比无所不能的智能手机，它可以用来拍摄，但是单反和专业摄像机依然存在；它可以用来照明，但是手电筒依然存在；它可以用来听歌，但是专业音响设备依然存在....但是使用单反、专业摄像机、手电、专业音响的人将越来越少，因为大部分人认为智能手机的功能对他已经够用了，从而选择了更便捷的“超融合”方案。

5、KV 数据库

一、KV 数据库分类

Google 在 2003 年至 2006 年间发布了一系列重要的论文，如 Google File System、Google Bigtable 和 Google MapReduce，奠定了分布式数据系统的基础。同时，Amazon 发布了 Dynamo 论文，第一次在非关系型数据库领域引入了数据库的底层特性，为后来的 NoSQL 数据库的发展奠定了基础。受此启发，NoSQL 数据库如 BigTable、HBase、Cassandra 等应运而生，这些 NoSQL 数据库解决了高并发读写、多结构化数据存储等问题，但其设计思路放弃严格的 ACID 事务处理、一致性和 SQL 查询。NoSQL 数据库有许多种，KV 数据库（Key-Value Database）便是其中一种常见的 NoSQL 数据库。

KV 数据库是一种存储和检索数据的数据库模型，它使用简单的键值对来组织和访问数据。随着时间的推移，KV 数据库经历了不断地发展和演进。早期的 KV 数据库主要关注于提供高性能的存储和检索功能。其中一些代表性的数据库包括 BigTable、MemcacheDB、LevelDB、Tokyo Cabinet、Tuple space 和 TreapDB 等。这些数据库通过使用哈希表、B 树等数据结构来实现快速的键值对查找和存储操作。它们在性能方面取得了显著的突破，并被广泛应用于各种场景，如缓存、会话存储和分布式存储等。

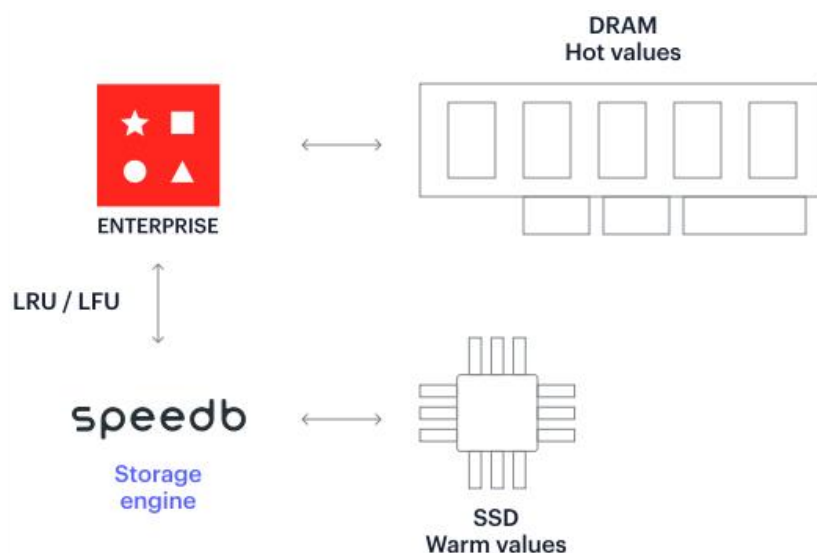
随着内存价格的下降和计算机硬件的发展，KV 数据库开始关注于提供基于 RAM 的存储引擎。这些引擎将数据存储在内存在中，以实现更高的读写性能。一些流行的基于 RAM 的 KV 数据库包括 Oracle Coherence、memcached、Citrusleaf database、Velocity 和 Redis 等。这些数据库通过将数据存储快速的内存介质中，实现了更低的访问延迟和更高的吞吐量。它们在需要极高性能和低延迟的应用中得到了广泛应用，如实时分析、高速缓存和会话管理等。

此外，一些 KV 数据库开始采用 Paxos 算法来实现数据的一致性和可靠性。Paxos 算法是一种用于分布式系统中达成共识的算法，它可以确保多个节点之间的数据一致性。Keyspace 是一个基于 Paxos 算法的 KV 存储系统，它提供了高可用性和强一致性的数据访问。通过使用 Paxos 算法，Keyspace 可以确保在节点故障或网络分区的情况下，数据仍然可靠地存储和访问。

二、KV 数据库 Redis

在 2009 年开源的内存数据库 Redis，作为当前最流行的 KV 数据库之一，以其丰富的数据结构、可扩展性（通过 Redis Module）和强大的生态系统，已经成为 KV

数据库事实标准。然而，作为内存数据库，Redis 在内存使用量超过一定阈值（如 16GiB）时，面临着一系列问题，包括内存容量限制、单线程阻塞、长时间启动恢复、内存硬件成本高、缓冲区易写满、一主多从切换代价大等。尽管 Redis 7 引入了多 I/O 线程功能以提升读写效率，但用户命令执行逻辑仍然是单线程处理，导致 CPU 开销增加而吞吐量没有显著提升。



为解决这些容量问题，Redis 企业版本引入了 Auto Tiering 功能（开源版本未提供此功能），该功能允许 Redis 数据库同时使用 DRAM 和 SSD 两种存储介质，以实现更高性能和可扩展性。Auto Tiering 将数据分为两类：热数据和冷数据。热数据是经常访问的数据，而冷数据是偶尔访问的数据。该功能将热数据存储于 DRAM 中，将冷数据存储于 SSD 中。根据其官方测试文档 Redis Enterprise Auto Tiering，使用多级存储后，吞吐量可提升 2 到 3 倍，但由于其单线程逻辑处理架构，适宜的数据容量上限仍未突破 100GiB。

为了突破 Redis 容量上限并解决内存不足的问题，许多人开始尝试在兼容 Redis 协议的前提下实现基于磁盘的 Redis 类数据库：

1. 2012 年开发的 SSDB，基于 LevelDB，兼容 Redis String/ZSet/Hashtable/List 数据结构。
2. 2013 年开发的 Ardb，基于 LevelDB/RocksDB/WiredTiger 等存储引擎，兼容 Redis 的多种数据结构 String/Hash/List/Sorted Set/Set/Geo/Bitmap/HyperLogLog/Stream，这个 DB 的一个创新点是实现了 GeoHash，后来被移植到 Redis 上。
3. 2015 年开发的 Pika，基于 RocksDB，兼容了 Redis 的 String/Hash/List/Zset/Set/Geo/Hyperloglog/Pubsub/Bitmap/Stream/ACL 等常

用的数据结构，能够缓存热数据实现冷热分级存储，单节点可以管理 x TiB 级别规模的数据量。

4. 2019 年开源的 Tendis，同样基于 RocksDB 实现，支持单机、主从、集群等多种部署模式，兼容 Redis 的绝大部分数据结构与命令。

三、KV 数据库演进

在云原生时代，基于磁盘的 KV 数据库的容量仍受到磁盘容量的限制，若不进行架构改造，可能导致资源浪费。

2012 年 Google 发布的 Spanner 论文，引入了 TrueTime 概念，创新性地将 SQL 和分布式事务引入第一代 NoSQL 系统，确保强一致性，使数据库发展进入了 NewSQL 阶段。开源实现如 CockroachDB 和 TiDB 取得了巨大成功，弥合了 NoSQL 和一致性之间的差距。后续的分布式 KV 存储大多支持分布式事务和数据强一致性，如 FoundationDB 和 TiKV，同时还有一些产品支持最终一致性，如 Tendis。一些 NewSQL 厂商奉行 “One Size Fits All” 理念，提供满足所有应用场景需求的数据库，例如 Oceanbase 的 OBKV 和华为的 GeminiDB Redis。

2014 年，AWS Aurora 提出了云原生数据库的 “日志即数据库” 理念，通过解耦存储与计算层，充分利用分布式文件系统的多副本和高可用特性，以及新型硬件加速技术，既保证了数据的强一致性又实现了极高的扩展性和性能表现，为 KV 数据库的未来发展提供了另一种演进路径。存算分离的理念包括弹性存储和 Serverless 数据库，极大地简化了数据库的管理和运维。一个典型例子是 Rockset 公司推出的基于 RocksDB 的 KV 数据库服务 RocksDB Cloud，通过与 S3 等云服务集成，提供了一种利用公有云存储服务进行持久化和分层存储的解决方案，使 KV 数据库可以更加灵活高效地部署于云端环境。

存算分离技术其本质思想就是通过读写缓存，做到存储成本和性能的动态平衡。计算节点写操作是把 WAL 日志写入写缓存，存储介质可以是本地磁盘或者 EBS 云盘抑或是消息队列，读缓存负责物化 WAL 日志为 Page 数据并存入对象存储。早期的存算分离被理解为计算和存储分离，最新的存算分离已经演化到计算、日志和存储三层分离架构，有的 DB 基于 RDMA 技术甚至做到了计算、内存、日志和存储等多级模块弹性伸缩。以 NeonDB 为例，其 SafeKeeper 负责存储 WAL 充当写缓存，PageServer 则负责重放 WAL 然后生成 Page 数据存入 S3，并缓存 Page 数据以提高读性能。

Serverless DB 一般会以一定量的 CPU\Memory\Network\Disk 为一个 Resource Unit 向租户售卖服务，然后对资源池中的 CPU\Network\Disk 资源进行压

缩进行超卖以保证超额利润，一般不会对内存进行压缩。超卖的前提当然是 DB 提供多租户服务，目前市面上的 Serverless DB 都会为每个租户提供一个 Proxy，以隔离后端资源并提供负载均衡、连接保持，多个租户可共用一个 Proxy。

综合分布式系统、NewSQL、Serverless 云原生的发展历程，2024 年 KV 数据库的云化趋势将更加明显：弹性存储、存算分离、多租户隔离等，将为 KV 数据库领域带来更灵活、可扩展且易管理的解决方案。在不断挑战容量规模的同时，数据库技术不断演进，为各类应用提供更为优越的性能和可用性。

6、多模态数据库

一、AI 时代的多模态数据库

（一）大模型带来的 AI 应用落地范式转移

自 2006 年深度学习时代以来，经过近 20 年的发展，越来越多的企业已经构建了成熟的 AI 应用开发和运维体系。这个体系通常包含三层结构：底层是机器学习平台，中间层是 AI 服务，顶层则是各种企业 AI 应用。在前-深度学习时代，AI 的实际应用开发和落地面临着多项挑战：数据方面，传统的小模型在泛化能力上存在局限，难以与企业的广泛私有数据相结合以提供外部服务，导致企业数据在面向 AI 的应用中未能实现有效连通和利用；在 AI 服务生态方面，尽管企业希望直接复用现有 AI 服务以快速构建应用，但除了少数服务（如 OCR、语音识别）实现了较高的可复用性外，大多数情况下企业仍需为特定数据和业务场景开发新算法和模型，影响了 AI 应用的落地效率。

然而，随着以 Large Language Models（LLM）大模型技术为核心的 AI 技术的迅猛发展，AI 正从前-深度学习时代迈入大模型时代，标志着 AI 时代的一个重要转折点。这种新的技术范式正在重塑企业 AI 应用的落地方式，并加速企业向全面智能化升级的步伐，同时也引领着 AI 应用开发和运维的新变革。大模型能够利用更多的企业数据进行智能化应用开发，不再局限于有限的数据集；其卓越的泛化能力使 LLM 大模型能够处理多种多模态的 AI 任务，从而节省了各种模型开发时间和多模型运维成本；此外，数百亿参数级别的大模型不仅具备强大的泛化能力，还展现出卓越的智能涌现能力，使得模型效果相比前深度学习模型有显著提升。

所有这些优势共同为 AI 在企业中的落地奏响了前进的号角。在实际应用方面，大模型技术的引入为各行各业提供了新的解决方案。例如，在金融服务行业中，大模型可

以帮助银行更精准地分析信贷风险，提高交易安全性；在医疗领域，它们能够协助医生进行更准确的疾病诊断和治疗方案的制定。此外，大模型还在客户服务、市场营销、供应链管理等多个领域展现了其强大的应用潜力。这种技术的广泛应用不仅提升了企业的运营效率，也为企业带来了更高的竞争优势。企业战略层面，大模型技术的引入正成为推动企业转型和创新的关键因素。通过结合大数据和先进的 AI 技术，企业能够更深入地洞察市场趋势，精确预测用户需求，从而做出更加明智的战略决策。这不仅有助于企业在激烈的市场竞争中保持领先地位，也为其长期发展奠定了坚实的基础。

（二）AI 应用的私有记忆体--多模态数据库

随着大模型技术的兴起，各企业正逐步尝试结合私有数据来推动 AI 的实际应用。在这一过程中，企业面临着一系列新挑战：

- **大模型与企业数据的结合：**大模型可能出现幻觉或不可解释性问题。企业如何确保所使用的 AI 技术能够准确理解和运用这些数据，同时满足对精准知识和可解释性的高要求，是一个重要考虑。
- **数据安全性：**在处理海量数据时，保护信息不被泄露是至关重要的。企业必须确保在整个数据处理和 AI 应用的过程中，所有的敏感信息都得到充分保护。
- **数据的专有性和个性化应用：**除了安全性问题，企业还需关注数据的专有性，即如何更深入地理解自身数据，并根据自己的特定环境和需求做出最适合的决策。这意味着 AI 技术不仅要能够处理通用数据，还应能够适应企业特定的业务场景和数据特点，实现真正的“个性化”智能解决方案。

为了应对这些挑战，企业需要构建更加智能和灵活的数据库系统，并对大模型的选择和应用进行细致考量。这不仅涉及技术层面的创新，还要求企业在策略和管理层面作出相应的调整。关键在于确保 AI 技术能够在尊重数据安全和保护个性化需求的基础上，为企业带来真正的价值。其中，构建适应大模型时代的多模态数据库是实现这一目标的关键基石。

在大模型时代，AI 应用的效能不仅取决于算法本身，还极度依赖于支撑这些算法的数据基础设施和数据库的强大与灵活性。因此，设计和实施一种能够支持传统结构化数据同时高效处理和存储各种非结构化和半结构化数据的数据库变得至关重要。这种数据库需要能够综合处理包括图、文本、向量等多种数据类型，以适应 AI 应用中日益复杂的需求。例如，在公安领域，为客户提供的 AI 应用需要实现更为复杂的多模态场景：根据嫌疑人的画像，查询特定时间内在城市哪些路口出现，同时了解与嫌疑人有关联的特定年龄段的男性信息。这要求使用多模态数据库，能够处理图片、视频、GIS 数据、

文本等多种数据格式。这种复杂数据的处理不仅对数据库的存储能力提出了挑战，还要求数据库具有高级的查询和分析功能，以便快速有效地提取和分析相关信息。

（三）以可解释性（图）为基础构建的 ArcNeural 多模态数据库

在当前的数据处理和数据库技术领域，我们正在经历一个重要的范式转变：从传统的存储和计算模式，转向更加强调智能化的记忆与推理。随着 AI 技术的深入发展，我们对数据的关注正在从单纯的数量积累和查询速度提升，转变为对数据整体的深入洞察和有效利用。

在这一转变过程中，可解释性已经成为客户的核心需求。虽然当前 AI 技术的主要进展依赖于概率系统，但我们相信，**符号系统的本质可解释性是对概率方法的重要补充**。在此背景下，图技术作为一种有效的符号系统，与基于概率的机器学习技术相结合，共同构成了可解释智能的基石。ArcNeural 多模态数据库正是基于这种综合视角构建，它不仅提供深度的数据分析能力，还以直观且易于理解的方式呈现复杂数据间的关系。

构建以图为基础的多模态智能引擎面临着几个关键挑战：首先，范式的转移从单纯的存储和计算，转向对全面数据理解的智能化推理；其次，图作为符号系统，与基于概率的机器学习相辅相成，是构建可解释智能的必要条件；此外，不同于传统的关系数据库，图结构不强调数据的本地性和局部性，而是更加重视数据间的关系；现实世界的数据不规则且多模态，要有效利用这些数据，必须适应它们的这些特性；最后，实时更新不仅涉及数据本身，还包括数据间关系的动态变化，这些都需要得到及时处理。

基于这样的理解，我们选择以图技术为核心来构建 ArcNeural 多模态数据库。ArcNeural 的设计重视数据的多样性和多模态特征，并强调数据及其关系的实时更新能力。这一特性使得数据库能够有效应对现实世界中数据的不规则性和持续变化。因此，ArcNeural 多模态数据库不仅是一个数据存储和处理的工具，它更是推动数据智能化应用的平台，提供了对数据的深入理解和高效推理的强大支持。

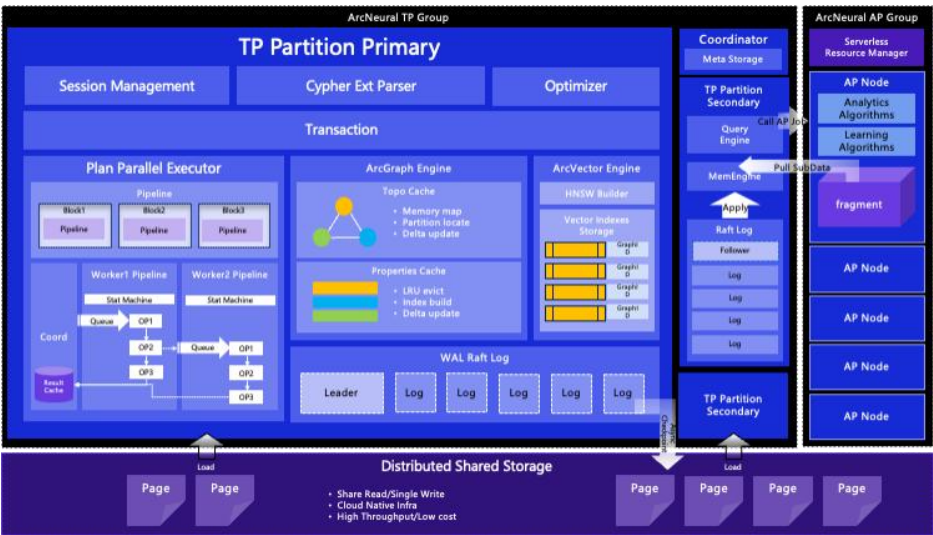
1. 架构设计

(1) 存储层：ArcNeural 多模态数据库的底层基础是其存储层。它采用了存储与计算分离的架构，这一设计充分体现了云原生环境的特性。这种架构不仅使得数据库能够灵活地适应用户现有的存储解决方案，同时也支持多种存储模式，从而增加系统的灵活性和扩展性。面对的主要挑战包括处理数据多样性以及满足实时更新的需求。

(2) TP（事务处理）计算节点：在设计 TP 计算节点时，我们并没有采用单一代码基础实现 HTAP（混合事务/分析处理），而是考虑到大多数图场景中，TP 和 AP（分析处理）通常有着截然不同的使用场景。在 TP 层面，ArcNeural 实现了一个完整的分布

式数据库系统，并支持多模态引擎，目前支持图和向量数据模式，并计划未来进行更进一步的扩展。

(3) AP（分析处理）部分的创新：AP 部分采取了 serverless 的形态，这意味着它不需要持续运行，而是根据需求临时在资源池中启动分析节点。这种设计使得数据能够实时从 TP 流向 AP，完成分析后还能将结果回传至 TP 节点进行数据更新等后续操作。这种机制使得 TP 和 AP 两套系统得以有机结合，各自发挥最大的优势。



ArcNeural 多模态数据库的设计旨在提供高度的灵活性和可靠性，同时优化性能。包括以下几个关键的技术考量：

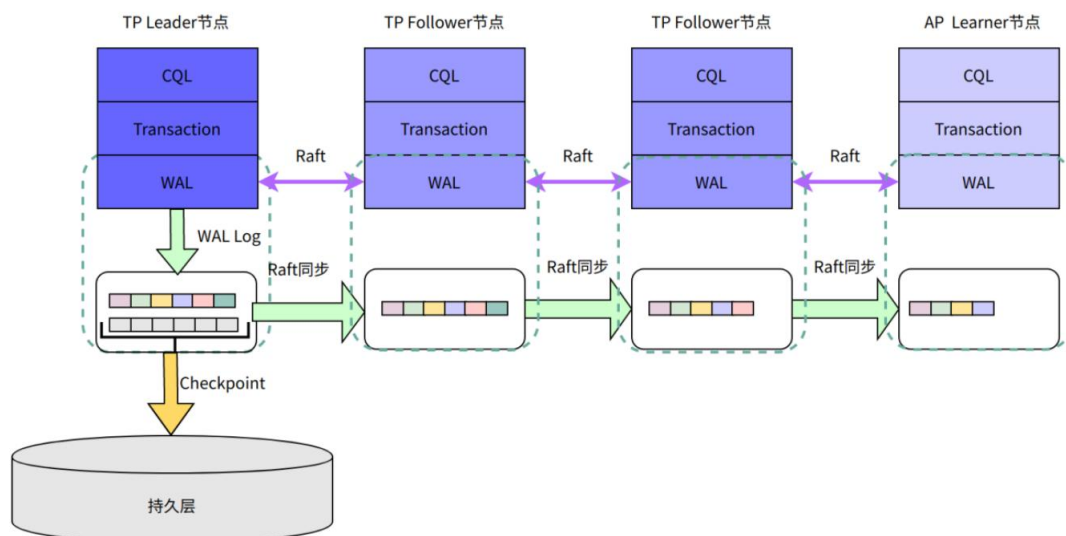
2.日志即数据

ArcNeural 沿袭了当前流行的“日志即数据”思想。存储节点只作为典型数据库的持久层，通过 Replay WAL 日志进行周期性 Checkpoint，存储节点可以随着系统数据量的大小进行动态扩缩容。同时根据业务需要可以对接单机 RocksDB，分布式 TiKV 或云厂商提供的对象存储系统（如 OSS）。

此外，ArcNeural 创新性地设计了计算层缓存 WAL 日志的方案，提出了 Semi-Stateful 概念。系统的 WAL 日志通过采用实现分布式一致性算法的 multi-raft 协议，来达到高可用和计算节点间的一致性。随着数据库的存算分离，传统数据库的业务 IO 瓶颈已经转化为计算层和存储层之间的 IO 瓶颈，在这种情况下，底层存储节点和网络时延对于响应时间有很大的影响，事务操作的性能也会很大程度上受到影响，尤其是分布式提交等复杂操作。因此，ArcNeural 在计算层增加本地高性能 SSD 日志存储，用以对 DML 的 WAL 日志进行存储，并定期对 WAL 日志进行 checkPoint，并写入底层存储引擎。这带来了如下优点：

- 在 DML 操作中，只需要记录 WAL 日志，从而显著减少了网络占用。

- 将复杂的 DML 操作由数据和日志的同步操作转变为连续的异步操作。
- 引入了一个分布式的、基于 Raft 协议同步的 WAL 日志服务，降低了网络层或数据存储层问题对数据库的影响。



ArcNeural 可采用灵活的分区策略，例如基于点的 ID 或点的 Type 进行 Hash，并采用边的分区跟随起始点的分区策略。通过这种方式，可以将图的点和边数据分为不同的 tablet 分区，每个分区都有多个 WAL 日志的副本，并使用 multi-raft 共识算法来实现 WAL 日志的分布式一致性。在同一个 Raft 组内，节点可以担任不同的角色，包括 Leader、Follower 和 Learner：

- **Group Leader**：负责接收分区 DML 操作所生成的 WAL 日志写入请求，并在 Raft 组内完成日志的同步及一致性确认，最后确保成功写入本地的日志 Store；
- **其它 Group Follower 节点**（一般是其他 TP 计算节点）：完成对 Raft Leader 新的日志进行同步确认，并在心跳超时后发起 Leader 选举；
- **Group Learner 节点**（通常是 AP 计算节点）：对最新的 WAL 日志进行同步，获取最新的 Delta 数据，并将其与从底层存储引擎获取的 Snapshot 数据进行合并，合并后的数据用来完成用户特定场景的 AP 计算（图计算）任务。

3.内存引擎设计

内存引擎层面，ArcNeural 不仅包含了完整的分布式数据库功能（如更新和事务处理），还专门为提升查询性能设计与底层存储引擎异构的内存存储。这包括图拓扑索引、属性索引以及向量索引，使得数据库在处理结构化、半结构化和非结构化数据时能够提供快速访问和高效处理。

存储和计算分离带来了一个优势，即计算节点无需配置很大的存储介质来存储图的点和边数据。这些数据可能非常大，而且大部分是冷数据（访问频率较低）。为了避免在每次查询数据时计算引擎都必须下探到存储节点，同时提高查询性能、降低计算节点和存储节点间的往来流量，ArcNeural 为每个计算节点提供数据缓存（Cache）层。其中包含：

(1)子图的拓扑结构, ArcNeural 目前采用邻接表数据结构来存储图拓扑;对于不可(常)修改的子图也会采用 CSR 来加速内存中多跳性能和矩阵计算性能;

(2)基于 LRU 的访问热点的点/边属性 Cache。

(3)index buffer 和 mvcc 多版本数据。

当前内存引擎在快速迭代研发和优化中，会借鉴业界领先的一些理念和设计，如跳表等高效数据结构，进一步提升内存引擎性能。

4.分布式查询

计算查询层采用 MPP 分布式架构，每个计算节点采用基于 Push 模式的向量化执行引擎，实现原理大致如下：

(1) **执行计划切割**：基于 QueryEngine 优化后的物理执行计划，执行引擎先结合算子的属性（例如是否是 pipeline breaker，如 Hash Join 算子），完成对执行计划进行切割，最终生成 DAG 图，基于 DAG 间的依赖关系，对其中的每个子 DAG 进行分布式并行调度执行；

(2) **计划分发**：根据执行计划中涉及数据 partition 分布信息，把执行计划分发到所涉及的 partition leader 节点；

(3) **分布式调度执行**：由 Co-ordinator 节点发起 pipeline，并在计算节点 culster 上分布式行执行；每个计算节点上的计算引擎采用 push 模型，并结合 Batch 方式进行向量化执行，同时算子内部利用协程进行计算加速（例如并行执行多节点的 Expand），高效完成执行计划的分布式并行执行计算；

在计算过程中，如果有些算子产生大量中间结果（如 aggregate 算子），计算引擎具备“Flush to Disk”的功能，可以物化部分结果到外存，从而避免 OOM 的问题。在有些场景中，如果计算引擎生成结果的速度快于 Client 拉取数据的速度，引擎会自动暂停 Pipeline 的执行，等待 Client 消费一定量 ResultCache 中数据后再恢复执行。针对图查询中典型多跳场景，根据节点的亲和性，对于需要到远端机器执行下一跳的节点集合，可以根据不同机器目的地进行缓存，并按批发送到远端机器。这些节点数

据发送到对端机器后，会启动新的子 pipeline 进行并行处理，然后根据 Plan 进行后续处理执行。

5.图的 TP（事务处理）和 AP（分析处理）融合

TP 和 AP 提供了统一的查询接口。在 TP 层面以图查询语言 Cypher 为基础，并参考了 ISO/GQL，针对 AP 能力进行了扩展，比如在语句中通过 MATCH Call 的方式调用 AP 的能力。当 TP 收到请求后，会进行语义分析来判断是否要使用 AP，如果需要则根据请求拆解算子，准备数据并调起 AP。

举一个简单的例子，假设我们想通过一条查询语句来了解一个社区中成员之间的 follower 关系，并通过这种关系来确定其中的意见领袖。我们可以撰写如下的查询语句

```
...  
  
MATCH (a:Person)-[e:follows]->(b:Person)  
  
WHERE a.community = 'ArcNeural '  
  
Call pagerank((a)-[e]->(b), 0.85, 10) YIELD *  
  
return id,result;  
  
...
```

TP 部分会进行处理，查询到谁 follow 谁，并得到一个子图，这是一个经典的数据库中的单跳查询。查询完成后，会调用 AP 系统。当 AP 系统收到这个命令后，它知道需要执行一个 PageRank 计算，就会自动从资源池中申请 AP 节点，组成一个小型的 AP 集群，然后将查询得到的子图发送过去进行计算。计算完成后，相关资源会被回收，实现 AP 请求粒度 Serverless。

最后，AP 引擎也可以独立运行来执行 T+1 这样的离线需求，同时还支持从外部存储比如对象存储或大数据系统里装载额外的数据，再通过我们上层业务平台提供的能力比如定时执行来实现自动化，这样就可以通过一套接口和一套系统来实现业务从 TP 到 AP 全场景的能力。

6.向量、文档和图的融合

在 ArcNeural 多模态数据库中，我们实现了向量、文档和图数据的高效融合。这种融合能力是基于图技术的强大表达能力，我们为图的属性增加了向量、文档和 JSON 的处理能力，从而支持复杂的多模态数据分析。

以电影知识图谱为例：电影和制作人作为图中的实体，形成了复杂的图拓扑关系。每部电影都有其独特的 ID 和元信息（如标题、年份、类型等）。除了存储这些基本信

息，ArcNeural 还支持将电影内容或影评的文字直接作为属性存储，或将其转化为向量数据后存储，从而支持更复杂的分析需求。

针对向量字段，我们通过 HNSW 索引来加速向量数据的检索。ArcNeural 在系统中加入了以下三个特性：

- 使用 SIMD 进行距离计算加速，与传统方式相比可以提升 4 倍性能。
- 支持属性过滤的向量检索，这不仅限于向量数据本身，还能结合其他字段进行综合查询。
- 利用图技术的基础，向量数据也获得了完整的数据库功能，如事务处理和统一查询。

我们可以来看一个查询的例子：想找到 2023 年上映的一部反欺诈题材的电影，并进一步查询该电影的导演还拍过哪些其他电影，这些电影里的主要演员有哪些。这个查询任务看起来比较复杂，但在 ArcNeural 中，用户可以使用一条简单的查询语句来完成。具体的查询语句如下所示。其中，黄色的部分是 WHERE 条件，用于筛选 2023 年上映且与反欺诈相关的电影。

```
Bash

MATCH

(m1:Movie)-[:DIRECTED]->(d:Director)<-[:DIRECTED]-(m2:Movie)-[:ACTED_IN]->(a:Actor)


WHERE m1.embedding<-> "[0.9,0.4, …… ,0.1]" > 0.9 AND jsonValue(m1.movie_info,
"$release_year") = "2023"


RETURN d.name, jsonValue(m2.movie_info, "$title") AS title, a.name;
```

其中包括两个部分，`m1.embedding<-> "[0.9, 0.4, …… , 0.1]" > 0.9` 用于计算向量接近度，在查询计划里是一个带 FILTER 的向量 INDEX SCAN，筛选出满足条件的电影；`jsonValue(m1.movie_info, "$release_year") = "2023"` 用于 Json 字段的过滤操作；然后继续向上传递，使用图算子对结果进行进一步的一跳、两跳查询，例如查询导演和他/她的其他作品。这种我们称之为 One Query 的设计确保了一种接口甚至是一条语句可以同时扫描多模态数据，可以极大地简化用户的业务研发复杂度。

（四）基于 ArcNeural 快速构建知识库智能助手

基于 ArcNeural 多模态数据库，我们开发了名为 Arc42 的企业知识库智能助手。Arc42 能够从 ArcNeural 中存储的多个内部数据源进行复杂的联合搜索，并提供精准

的答案。这些内部数据源主要包括技术文档、代码和组织架构等，使得 Arc42 能够回答基于内部知识的问题，例如“我有 V2.0 Compiler 设计相关的问题，该请教谁？”或“谁最懂大模型”等。Arc42 的用户界面简洁直观，可以轻松接入到飞书或钉钉等平台，作为聊天机器人使用。



除了 ArcNeural 多模态数据库的支持，Arc42 主要包括数据生产阶段和用户查询处理阶段：

1.数据生产阶段：Arc42 将内部数据源，包括技术文档、代码、组织架构等，构建多个图谱，以及与之相关的属性和文本片段。

2.用户查询处理阶段：

- (1) 接收用户的原始问题 q ;
- (2) 利用大型语言模型 (LLM) 和图谱的结构信息，将用户问题解析成符合图结构定义的意图链;
- (3) 将意图链分解为子计划，并依次执行查询和推理任务。

在使用 LLM 进行用户意图拆解的过程中，我们将访问图/向量数据库的不同 API 抽象为多种算子，并向 LLM 提供这些算子的功能描述、输入输出等信息。基于这些算子，我们定义了部分公式化的表达作为(Chain of Thought)COT 示例，使得 LLM 能够对不同的用户查询进行功能拆解，并生成查询计划。

例如，在一个典型的交互场景中，用户可能会问：“我可以把 Mahinda Perera 介绍给 Carmen Lepland 吗？”在这个例子中，Arc42 的用户意图识别模块将解析这个问题，并生成相应的执行子计划。下面的示例展现了 Arc42 处理查询的独特方法，结合了符号化风格的意图识别和基于 Cypher 风格的执行子计划。

(1) 用户查询：“Can I introduce Mahinda Perera to Carmen Lepland?”

(2) 用户意图识别模块输出：

- (Mahinda Perera, knows, Carmen Lepland)&both edge&(person, knows, person)
- (Mahinda Perera, hasInterest, symbol:t)&edge&(person, hasInterest, tag); (Carmen Lepland, hasInterest, symbol:t)&edge&(person, hasInterest, tag)
- 这种符号化的输出表示 Arc42 将查询转化为图查询和属性查询的组合。

(3) 执行模块生成的执行子计划：

- 查询：
 - ① (Mahinda Perera:person) - [:knows] - (Carmen Lepland:person) -> 回答 1
 - ② (Mahinda Perera:person) - [:hasInterest] -> (t: tag); (Carmen Lepland:person) - [:hasInterest] -> (t: tag) -> 回答 2

(4) 推理：

将答案、摘要和最终答案结合逻辑进行汇总。

(5) 解释输出：

将推理过程和原因以人类可读的方式输出。

这种结合符号化意图识别和执行子计划的方法允许 Arc42 处理相对复杂的用户查询。最后，执行模块将查询和推理的结果转化为自然语言，以直接和易于理解的方式生成最终回答。Arc42 通过这种方式不仅使得数据查询变得更加智能化和高效，还极大地简化了用户的业务研发复杂度。

这样，我们不仅实现了对复杂查询的快速响应，还提供了一种有效的方式来管理和利用企业内部的丰富知识资源。这个智能助手证明了 ArcNeural 多模态数据库在实际应用中的强大能力，同时也展示了智能化的知识管理和查询系统在当今数据驱动的时代巨大潜力。

二、面向 IoT 的多模数据库

与多模态数据库定义不同，多模数据库是指能够有效处理和管理多种数据模型的数据库系统，例如关系型数据、时序数据、文档型数据、图数据等多种形式的数据库。中国信通院在《数据库发展研究报告（2022 年）》中将多模数据管理列为九大数据库关键技术之一，中国通信标准化协会大数据技术标准推进委员会(CCSA TC601) 在《数据库发展研究报告（2023 年）》中也指出，多模数据库将逐步形成新的规范和使用方式。这种数据库的灵活性使其能够适应不同类型和结构的数据，特别是在物联网（IoT）技术和应用不断演进的背景下，面向 IoT 的多模数据库为应对物联网环境中产生的多样化数据提供了有效的解决方案，将成为推动物联网技术应用发展的核心支撑。

在迈入全新物联网时代的今天,面向 IoT 的多模数据库在物联网场景中有以下特点：

（一）以时序数据为代表的多样化数据整合

在智能家居、智慧交通、智慧城市、工业互联网等典型的 IoT 场景中，各类传感器、设备产生的数据类型多种多样，其中最为典型的是各类传感器的时序指标数据，此外还有结构化的关系型数据，非结构化的文本数据、图像和视频等。面向 IoT 的多模数据库可以高效整合这些多样化的数据，实现在一个数据库系统中对各种数据模型的存储和处理支持，既能保障数据一致性、传输转换等关键性能不受影响，还能大大降低开发、运维及管理成本。2023 年，浪潮 KaiwuDB 推出业内首款“分布式多模”数据库 KaiwuDB，支持时序数据、结构化、非结构化数据在同一库中归集处理，为用户提供更全面的信息和更统一的数据管理方式。

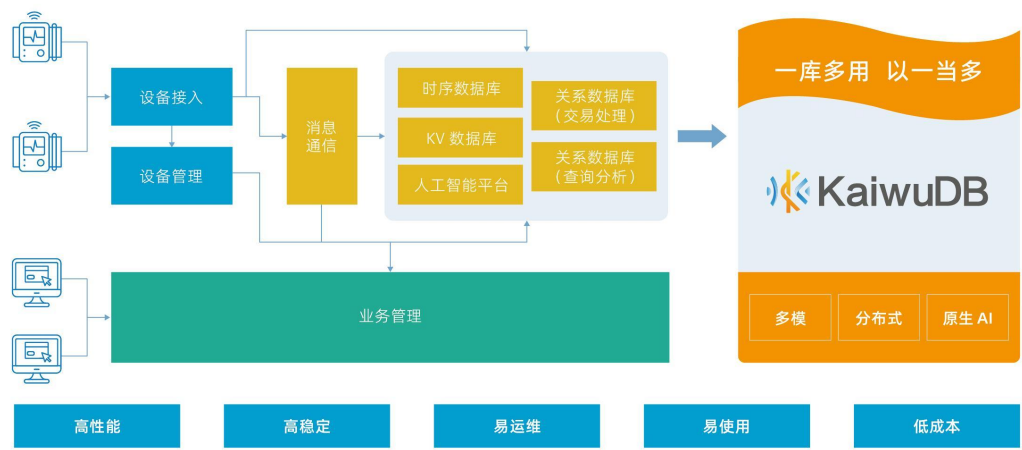


图 4：面向 IoT 的多模数据库（以 KaiwuDB 为例）

（二）实时性和低延迟

IoT 应用通常对数据的实时性和低延迟有极高的要求，尤其在智慧电网、工业互联网等领域。面向 IoT 的多模数据库通过对不同数据模型的写入、存储和检索进行有针对性地设计和优化，能够满足 IoT 环境下对实时数据处理和低延迟的需求，提高系统的响应速度。例如，浪潮 KaiwuDB 针对 IoT 场景下的时序数据，通过就地计算等技术支持百万级秒级写入和毫秒级精度数据写入，支持数据实时分析，千万笔数据聚合查询毫秒级响应。

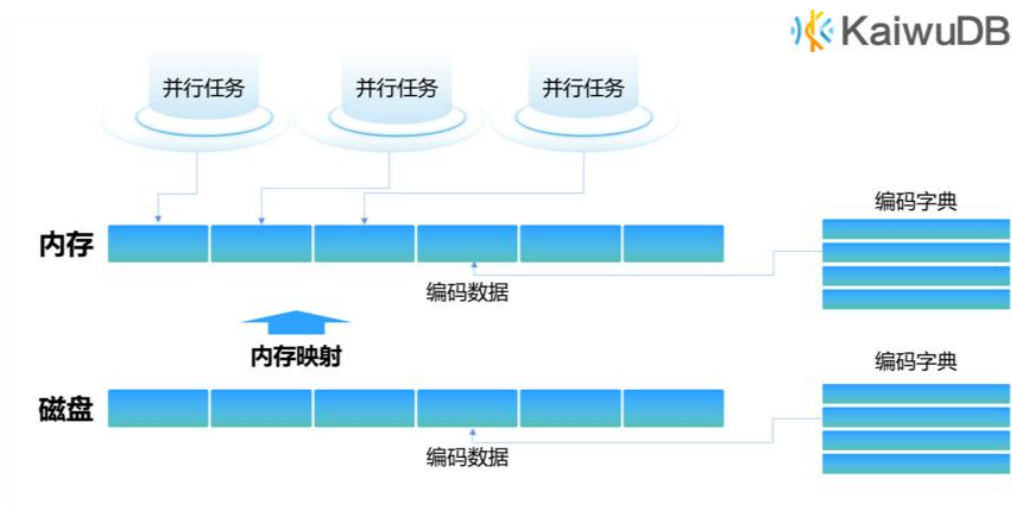


图 5: KaiwuDB 就地计算

（三）连接性和分布式架构

IoT 场景涉及大量连接的设备，而这些设备可能使用不同的通信协议和数据格式。面向 IoT 的多模数据库需要具备强大的设备管理功能，支持各种连接性需求，确保与不同类型的设备无缝集成和通信。同时由于产生数据源的设备通常基数大、分布广，面向 IoT 的多模数据库还支持“云-边-端”的分布式数据库架构，支持端侧轻量化部署；具备集群部署、数据同步、数据订阅等能力。

（四）智能分析和预测决策

IoT 生成的海量数据需要经过智能分析，以提取有价值的信息。除了通过就地计算、流式计算等技术实现多层级、多维度数据下钻，从海量数据中挖掘提取数据的价值，多模数据库还可以集成预测分析引擎，通过人工智能算法提供原生 AI 能力，从而提供更智能的数据分析方法，支撑用户科学管理决策。

总结：

综合而言，面向 IoT 的多模数据库不仅仅是数据库技术的延伸，更是适应物联网时代需求的关键工具。它为物联网应用提供了强大的数据管理和分析能力，推动了物联

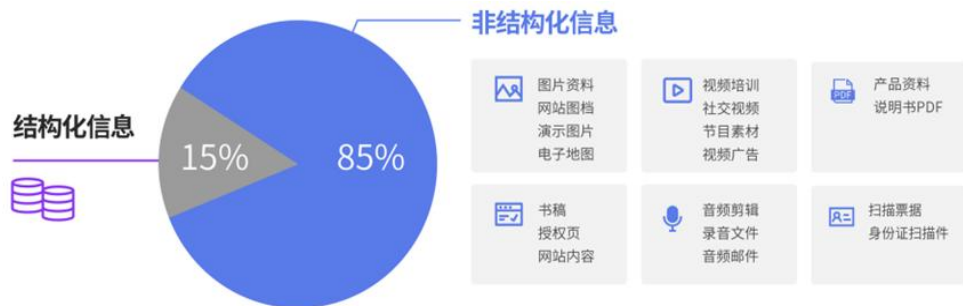
网技术在各行各业的广泛应用。未来，随着物联网技术的不断发展，面向 IoT 的多模数据库将在推动数字化转型、智能城市建设等领域发挥越来越重要的作用。

7、搜索型数据库

一、概述

随着数字科技的飞速发展和信息量的爆炸性增长，搜索引擎已成为我们获取信息的首选途径之一，典型的代表厂商如 Google。然而，随着用户需求的不断演变，传统的搜索技术已经无法满足人们对信息的实时性、个性化和多样性的需求。

在企业内部，这种需求更加显著。随着企业数字化转型的持续深化，非结构化数据正日益成为各类组织数据增长的主要来源，也是数据体系中至关重要的组成部分，蕴含着巨大的价值。如何高效地存储和利用非结构化数据的重要性也日益凸显。企业需要更高效地管理和检索内部的海量数据，以支持业务决策和运营需求。



据 IDC 数据预计，到 2025 年，80%的数据将是非结构化数据；而根据 Gartner 的数据显示，从 2019 年到 2024 年，非结构化数据容量预计将增加两倍。然而，目前非结构化数据面临着表现形式多样、管理复杂性高、价值挖掘难度大等诸多挑战。传统的数据库系统往往无法满足企业对实时性和多样性的搜索需求，为了解决这些挑战，以自动分词、倒排索引、相关度计算、向量检索引擎等技术为核心构建的搜索型数据库应运而生。这些数据库自上世纪 90 年代诞生以来不断发展演进，正在成为数据库领域中不可或缺的一个重要分支。

二、什么是搜索型数据库？

搜索型数据库早期又称全文数据库，或者企业搜索引擎，是一种专门用于存储和管理大规模文本数据，并支持高效的文本搜索和信息检索的数据库系统，不过随着技术不断发展和应用场景日益丰富，目前搜索型数据库不仅仅可以处理长文本数据，也可以处理常见的数值、日期等结构化数据，IP、地理位置信息、图片、音视频等非结构化数据，搜索型数据库的应用范畴不断拓展，正在由支撑业务系统检索加速、IT 运维可观测性、聚合查询分析等向多场景、多模态数据搜索方向发展。

典型的搜索数据库一般具有以下特点：

1. **灵活的索引能力**：搜索数据库能够处理多种类型的数据，包括文本、图像、音频、视频等非结构化数据。它们采用自动分词、倒排索引等技术，能够高效地处理不同格式和类型的数据，提供灵活的搜索和检索功能。
2. **高效的查询性能**：搜索数据库具有高效的查询处理能力，能够快速索引和检索大规模的数据。借助优化的索引结构和查询算法，搜索数据库能够在短时间内准确地返回与查询相关的结果，提高用户的搜索效率，常用于解决关系型数据库的高并发检索需求。
3. **支持复杂的搜索功能**：搜索数据库提供多样化的搜索功能，包括全文检索、模糊搜索、精确搜索、范围搜索、向量搜索、地理信息检索等。用户可以根据不同的需求和场景，灵活地选择和组合不同的搜索功能，以获取符合期望的搜索结果。
4. **高性能和可扩展性**：搜索数据库具有高性能和可扩展性的特点，能够处理大规模数据和高并发访问。它们采用分布式架构和并行计算技术，实现了水平扩展，能够满足不断增长的数据量和用户访问量的需求。

综上所述，搜索数据库具有处理非结构化数据、实时搜索和更新、多样化的搜索功能、个性化推荐和智能搜索、高性能和可扩展性、全面的搜索结果展示等特点，是处理大规模数据和提供高效搜索服务的重要工具。

三、搜索型数据库的应用场景



搜索型数据库在各行各业都有广泛地应用，以下是一些典型的应用场景：

1. **零售和电商**：在零售和电商行业，搜索型数据库被广泛应用于产品搜索和推荐系统中。通过搜索功能，顾客可以轻松查找所需商品，而个性化推荐系统则可以根据用户的搜索历史和行为习惯推荐相关的产品，提高购物体验 and 交易转化率。

2. **医疗保健**：在医疗保健行业，搜索型数据库被用于医学文献检索、疾病诊断和药物搜索等方面。医生和研究人员可以利用搜索功能找到相关的医学文献和研究成果，帮助诊断疾病和制定治疗方案。
3. **金融服务**：在金融服务行业，搜索型数据库被用于金融数据检索、市场分析和投资决策等方面。投资者可以通过搜索功能查找相关的金融数据和市场资讯，帮助他们做出更加准确的投资决策。
4. **制造业**：在制造业中，搜索型数据库被用于生产过程监控、质量控制和故障诊断等方面。工程师可以利用搜索功能查找相关的生产数据和技术资料，帮助他们解决生产中的问题和挑战。
5. **媒体和娱乐**：在媒体和娱乐行业，搜索型数据库被用于内容检索、版权管理和用户推荐等方面。用户可以通过搜索功能查找感兴趣的新闻、音乐和视频等内容，而个性化推荐系统则可以根据用户的搜索历史和偏好推荐相关的内容。
6. **教育和培训**：在教育和培训行业，搜索型数据库被用于学习资源检索、课程管理和学习分析等方面。学生和教师可以利用搜索功能查找相关的学习资源和课程内容，而学习分析系统则可以分析学生的搜索行为和学习表现，为教学提供参考和支持。
7. **IT 运维可观测性**：通过搜索型数据库，可以实时监控系统的运行状况、性能指标和日志数据，帮助运维团队及时发现和解决系统故障、性能问题和异常情况，确保系统的稳定运行。
8. **安全监测和威胁监测**：利用搜索型数据库对系统的安全日志进行审计和监控，监测用户的访问行为和系统操作，及时发现异常行为和安全事件。同时，搜索型数据库还可以与威胁情报数据集成，对内部日志数据进行关联分析，快速识别并应对各种安全威胁和攻击行为，保障系统和数据的安全。

综上所述，搜索型数据库在各行各业都发挥着重要作用，数据规模从 GB 到 PB 不等，体现在生活中的方方面面，为用户提供了高效、准确和个性化的信息搜索和检索服务，推动了各行业的发展和进步。随着搜索技术的不断创新和发展，搜索型数据库在行业中的应用将会越来越广泛，并持续为用户带来更加便捷和智能的搜索体验。

四、搜索型数据库的发展历程

搜索型数据库的发展历程可以概括如下四个阶段：



1. **起步阶段（1990 年代）：** 搜索数据库的雏形开始于上世纪 90 年代，当时以全文检索为主要技术手段，最初用于文档检索和网络搜索。典型代表包括 AltaVista、Excite 等。
2. **技术突破（2000 年代）：** 随着互联网的快速发展，搜索数据库开始应用于更多领域，如电子商务、社交网络等。Lucene、Sphinx 等开源搜索引擎的出现推动了搜索技术的进步。
3. **商业化发展（2010 年代）：** 搜索数据库进入商业化阶段，以 Elasticsearch 等为代表的商业搜索引擎崭露头角。企业开始大规模应用搜索数据库来管理和检索大量数据。
4. **智能化转型（2020 年代）：** 随着人工智能技术的发展，搜索数据库逐渐向智能化转型，开始引入机器学习、自然语言处理等技术，提供个性化推荐和智能搜索服务。同时，搜索数据库也在更多领域得到应用，如医疗保健、金融服务等。

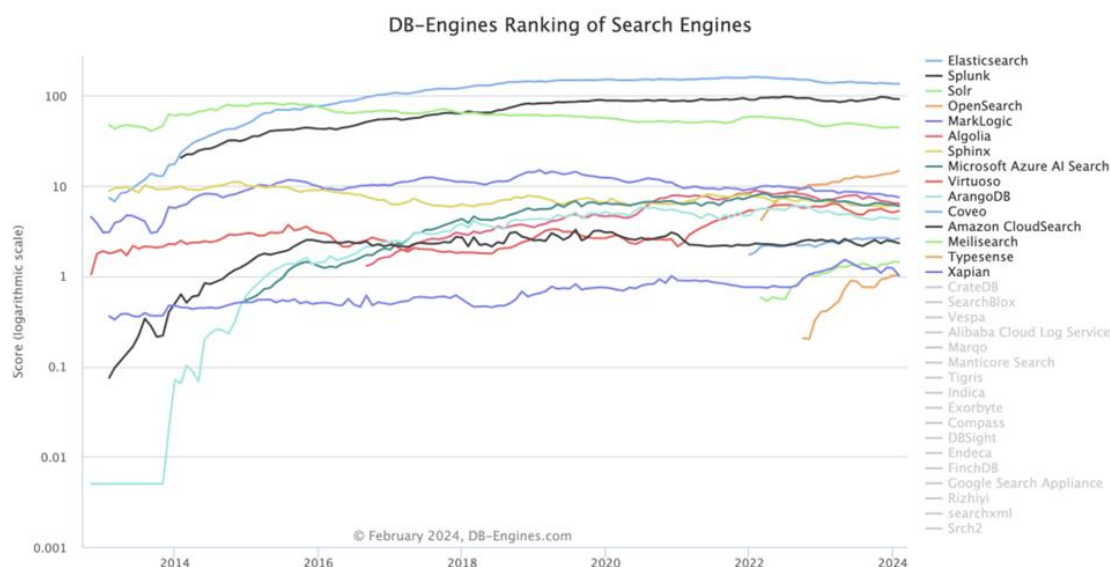
综上所述，搜索数据库经历了从起步阶段到技术突破、商业化发展再到智能化转型的发展历程，表明了其在信息检索领域的重要性和不断演进的趋势，并不推动着搜索技术的进步和应用范围的扩展。随着人工智能技术的不断成熟，搜索数据库将会在智能化、个性化等方面取得更大的进步，为用户提供更加优质的搜索体验。

五、搜索型数据库的发展情况

搜索型数据库市场上已经有不少成熟的产品和厂商，但是总的来说，搜索型数据库的界限范围有点模糊，当然其他数据库也有同样的问题，有很多数据库既是文档数据库，又是多模态数据库，还是向量数据库等等，而常见的搜索型数据库主要诞生于：

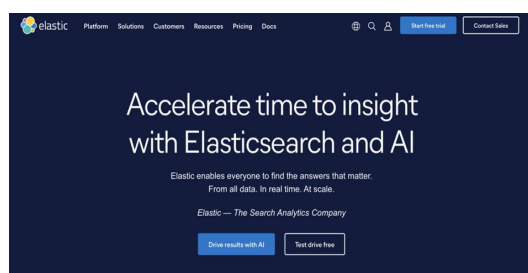
1. 由搜索引擎内核库发展而来的搜索数据库，如 Elasticsearch
2. 由其他数据库扩展而来的搜索数据库，如 Postgres Full-Text Search
3. 从零开始整体设计的搜索数据库：如 INFINI Pizza

通过流行的 DB-Engines 的搜索引擎排行榜，可以初探国外主流的搜索型数据库的流行趋势，如下图：



DB-Engines 的最新搜索引擎排名

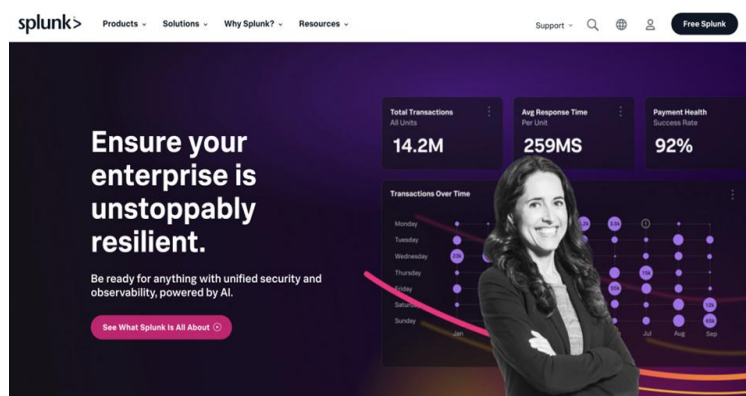
可以看到 Elastic 公司的 Elasticsearch 还是依旧保持强悍, 自从 Elasticsearch 十多年前掀翻了 Splunk 的桌子, 硬生生地在日志领域杀出一条新路, 随后大杀四方, 碾压整个搜索行业, 霸榜至今。Elastic 商业化增长稳健, 2023 年收入超过 10 亿美金。



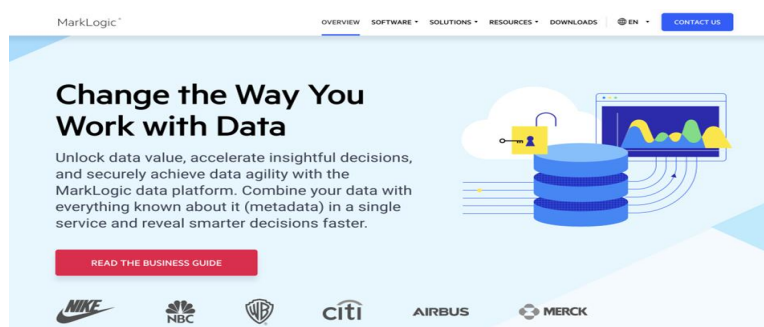
OpenSearch 是由 AWS 发起的 Elasticsearch 开源分支, 起因是由于 Elastic 针对云厂商采取的协议变更为 Elastic+SSPL, OpenSearch 基于 Apache 2.0 协议的 Elasticsearch 7.10 版本衍生而来, 目前也具备了一定的用户基础。



Splunk 是一款用于搜索、监控和分析大规模机器生成的数据的软件平台，主要用于日志和安全分析领域，属于商业闭源产品。2023 年中被思科（Cisco）以 230 亿美元现金收购，瞬间刷爆朋友圈。另外有意思的是，前四名除了 Splunk，底层都是 Lucene 内核。



MarkLogic 成立于 2001 年，自我定位是一个 NoSQL 多模态数据库厂商，也是商业闭源软件，生态成熟但是系统过于复杂，学习曲线较陡，2023 年初被 Progress Software 以 3.55 亿美元收购算是一个比较好的结局。



当然了，除了榜上的这些产品，还有很多优秀的挑战者正摩拳擦掌，跃跃欲试。如下面的这些项目：

vespa、Rockset、Doris、Clickhouse、quickwit、Pinot、SingleStore、qdrant、milvus、algolia、meilisearch、typesense、Manticore Search 等等。这些项目不一定是自己定位是搜索型数据库，有侧重在 AI 领域的，有侧重在实时分析领域的等等，可谓各有千秋，不过都具备一定的搜索和分析能力，不出意外，基本上每家都要号称吊打 Elasticsearch 一番。

六、国内搜索型数据库的发展情况

搜索型数据库已经成为企业事实上的重要基础设施，而国内搜索型数据库的发展近年也是开始得到重视，2023 年初，由中国信通院云计算与大数据研究所牵头，依托

中国通信标准化协会大数据技术标准推进委员会，联合拓尔思、极限科技、星环科技等 30 余家企业编制的《搜索型数据库技术要求》正式出炉，该标准已成为行业内搜索型数据库技术选型和产品开发的风向标，极限科技的 INFINI Easysearch 率先通过了该标准。

墨天轮社区也开辟了搜索型数据库的排行榜，共有 7 家企业的产品上榜：

排行	上月	半年前	名称	模型	数据处理	部署方式	商业模式	专利	论文	案例	资质	书籍	得分
1	↓↓↓ 33	↓↓↓ 30	Easysearch	搜索	-	☒	☑	2	0	0	3	0	33.93
2	↓↓↓ 44	↓↓↓ 49	Beaver	搜索	-	-	☑	36	0	0	5	0	18.40
3	↑↑↑ 282	↑↑↑ 282	Havenask	搜索	-	-	-	0	0	0	0	0	2.07
4	↓ 246	↓↓↓ 214	Scope	搜索	-	☒	☑	0	0	0	0	0	0.04
5	↓↓↓ 247	↓↓↓ 225	TRS Hybase	搜索	-	-	☑	0	0	0	0	0	0.03
6	↓↓↓ 269	↓↓↓ 260	Tera	搜索	-	☒	☑ ☒	0	0	0	0	0	0.00
7	↓↓↓ 276	↓↓↓ 226	思极亿搜	搜索	-	-	☑	0	0	0	0	0	0.00

墨天轮排行榜-搜索型数据库排名（2024 年 5 月）

国内搜索型数据库的市场还在起步阶段，厂商和可选的产品也还比较少，不过随着市场的成熟，相信未来将迎来一波高速的发展。

七、搜索型数据库的趋势前瞻

技术在演变，场景在演变，数据也在演变，搜索数据库领域的发展也呈现出多个显著的趋势，这些趋势将进一步推动搜索技术的演进和应用范围的扩展。笔者观测到的主要发展趋势包括以下方向供参考：

（一）趋势一：实时搜索与分析

- 1. 实时搜索是搜索数据库领域的一个重要发展趋势，业务应用都在朝实时方向演进，用户对信息的即时性需求不断增加，要求搜索结果能够及时反映最新的数据和内容。
- 2. 实时搜索技术通过实时索引和实时更新机制，能够实现快速的数据检索和更新，提供与时俱进的搜索结果，满足用户对信息的即时性需求。
- 3. 目前以 Lucene 为内核的搜索型数据库基本上都只能做到 NRT（近实时）搜索，并且频繁更新带来的挑战和资源的浪费比较高，如果能做到更高效的实时性，可以大幅提升用户的搜索体验和实时决策能力。

（二）趋势二：多模态混合搜索

1. 多模态搜索是指在搜索过程中同时考虑多种信息形式，如文本、图像、视频等，以提高搜索结果的准确性和全面性。
2. 这种技术能够通过分析和理解多种信息形式之间的关联性，为用户提供更加全面、丰富的搜索结果，适用于需要综合不同媒体形式的搜索场景。
3. 现实世界的数据越来越复杂化，非结构化数据的利用的场景也越来越多，多模态可以为业务提供更加灵活的分析和探索能力，混合搜索的能力非常具有吸引力。

（三）趋势三：AI 智能语义搜索

1. 大模型、AI 智能搜索技术的探索可谓是一日千里，通过利用人工智能技术来实现搜索过程中的智能化、语义化和个性化，结合自然语言处理、机器学习等技术分析用户意图，提供更加智能、个性化的搜索服务。
2. 随着大模型的兴起，搜索数据库开始采用像 RAG（Retriever-Reader for Generative Question Answering）这样的大型预训练模型来提升搜索的效果。RAG 模型结合了检索器和阅读器的功能，能够实现更加准确和全面的搜索结果，为用户提供更加智能和个性化的搜索服务。
3. 搜索型数据库可谓是 AI 落地最好的试验田，Elasticsearch 通过拥抱 AI 和大模型，目前股价又重回巅峰，可喜可贺。

（四）趋势四：云原生、存算分离、Serverless

1. 随着云计算技术的发展，搜索数据库正逐渐向云原生架构转变。云原生搜索数据库利用容器化、微服务架构等技术，实现了更高的灵活性、可扩展性和容错性，为企业提供了更加稳定和高效的搜索服务，并且成本更低，更加弹性。
2. 存算分离是搜索数据库发展的另一重要趋势。通过将存储与计算分离，搜索数据库可以更好地适应数据存储和计算需求的变化，提高系统的性能和效率。存算分离技术使得搜索数据库能够实现更高的并发访问和更快的数据处理速度，为用户提供更加流畅和稳定的搜索体验。
3. Serverless 提供开箱即用的体验，成本更低，使用更加灵活，也是目前很多搜索服务提供商正在积极探索的方向。

（五）趋势五：增强现实搜索

随着增强现实技术的发展，尤其是 Apple 发布的头戴式 Vision Pro，一部革命性的空间运算设备，将数位内容无缝融入实体世界，而搜索技术也将逐渐与增强现实相结合

合，为用户提供更加直观和沉浸式的搜索体验。增强现实搜索能够将搜索结果与现实世界相结合，结合 AI 技术为用户提供更加个性化和便捷的搜索服务，这是一个全新的领域，也意味着巨大的机会。

（六）趋势六：现代硬件的高效利用

1. 现代硬件及软件运行环境已发生翻天覆地的变化，片上计算，边缘计算，FPGA，DPU，GPU，一台设备几百核上 TB 内存已经成为现实，可运行之上的软件却还是停留在几十年前的架构。如 Elasticsearch 其核心 Lucene（及类似实现）是在 1997 建立的，距今已有 27 年了，虽然也在与时俱进，但是部分架构和设计理念已不具备先进性。
2. 在现代的硬件上采用更先进的算法，更新的数据结构、更新的设计理论，利用最新的 CPU 指令集，向量化，批处理，充分发挥多核、大内存和 SSD 的优势，从而达到更高的效率，更低的成本，去解决之前不可能实现的问题，大有可为，也是下一代引擎需要关注的方向。



随着各类数据库功能的边界越来越模糊，应用场景高度交叉重叠，市场竞争也变得白热化，不过笔者认为垂直领域的搜索型数据库机会还是很大，而想做大而全的数据库产品已经没有太多的市场生存空间，一定要在垂直领域有特别专注的地方，INFINI Labs 正在基于 Rust 研发的下一代搜索引擎 INFINI Pizza，就侧重于面向终端用户场景，解决海量数据更新情况下，同时满足高并发和低延迟的核心业务实时检索需求。

总结：

综上所述，搜索数据库领域正处于快速发展的阶段。随着互联网数据量的不断增长和用户需求的不断变化，搜索数据库技术将不断创新和进步，以满足用户对信息获取的更加即时、个性化和多样化的需求。未来，随着人工智能技术的进一步发展和应用，搜索数据库将会变得更加智能化、普及化和多样化，为用户提供更加高效、准确和个性化的搜索服务，推动互联网信息的更加便捷获取和利用。

8、图数据库

一、分布式原生图数据库、图计算、图数据中台技术生态

（一）图数据库的技术发展

图数据库是一种以图结构存储和管理数据的数据库，它能够有效地表示和处理具有复杂关系的数据。近年来，图数据库技术发展迅速，随着大数据技术和人工智能的快速发展，图数据库作为一种高效管理复杂关系数据的数据库系统，逐渐成为数据管理领域的重要技术之一。图数据库利用图论的理论，以图的形式存储实体之间的关系，相比传统的关系型数据库，在处理复杂关系和深度查询时表现出显著的优势。

1. 分布式原生图数据库的特点

分布式原生图数据库是设计来存储和管理大规模图数据的数据库系统。它们的设计原则是从底层支持图数据的分布式存储和计算，以实现高性能、高可扩展性和高可用性。基于分布式技术，原生图数据库实现几个关键特点包括：

- (1) 高性能：通过优化数据模型和查询算法，实现快速地图遍历和复杂查询处理。
- (2) 高可扩展性：支持水平扩展，可以通过增加更多的服务器节点来扩大数据库的容量和提高处理能力。
- (3) 强一致性：即使在分布式环境中，也能保证数据的准确性和一致性。
- (4) 容错性与高可用性：通过数据复制和故障转移机制，确保系统在面临硬件故障或网络问题时仍能正常运行。

图原生存储作为图数据库的核心技术之一，近年来在学术界和产业界引起了广泛关注。其针对图数据特点进行的优化设计，能够有效提升图数据库的存储和查询性能，为大数据时代的海量图数据管理提供了强有力的技术支撑。

在存储方面，图原生数据库采用了免索引邻接技术，这种技术通过在节点内部直接存储指向相邻节点的指针，使得遍历操作极为高效，因为它消除了查找邻接节点时访问全局索引的需要。然而，为了支持更复杂的查询和操作，图原生存储系统通常还会设计专门的辅助索引架构，如主键和属性索引。在查询优化方面，技术上也有很多路径进行优化，如并行查询和谓词下推等查询优化技术，并行查询通过分布式计算来加速大规模图数据的查询处理，谓词下推技术则在查询执行早期过滤掉不符合条件的数据，减少后续处理的数据量。

未来图数据库的发展趋势将显著受到硬件加速和对更多数据类型的支持这两大技术点的推动。随着图数据的规模和复杂性不断增加，传统的 CPU 处理能力已经难以满足对高效图数据处理的需求。因此，硬件加速，特别是利用 GPU（图形处理单元）和 FPGA（现场可编程门阵列）成为提升图数据库查询性能和图算法执行速度的重要手段。GPU 的并行处理能力使其非常适合进行大规模图数据的并行计算，而 FPGA 的可编程性则为针对特定图处理任务的定制化硬件加速提供了可能。

同时，为了更全面地理解和分析复杂的现实世界关系，图数据库正向着支持更多数据类型方向发展。这包括但不限于时序数据、空间坐标、多媒体数据等，这些数据类型的引入将极大地丰富图数据库的应用场景，比如在社交网络分析、地理信息系统、复杂系统建模等领域的应用将变得更加高效和深入。通过整合这些多样化的数据类型，图数据库能够提供更为细致和全面的分析结果，帮助用户从多维度理解数据间的复杂关系，进而做出更为准确的决策。这两大技术点的发展预示着图数据库在未来将变得更加强大和灵活，为处理和分析复杂数据关系提供了更为高效和全面的解决方案。

2.图数据库查询语言标准发展

图数据库查询语言的发展正处于一个关键的转折点，特别是随着图查询语言标准化进程的推进。目前，图数据库领域存在多种查询语言，如 Cypher、Gremlin 和厂家自定义变种等，它们分别由不同的图数据库系统所采用。这种多样性虽然丰富了图数据库的生态系统，但也带来了学习成本的增加和跨系统兼容性的问题。为了解决这些问题，行业内正努力推动图查询语言的标准化，其中最具代表性的努力是 GQL（Graph Query Language）的提出和发展。

GQL 是一种旨在成为图查询语言国际标准的新语言，由 ISO/IEC 在近年启动并推动。GQL 的设计综合了现有图查询语言的优点，目标是提供一种强大、灵活且易于学习的标准化图查询语言。它不仅支持高效的图数据查询和更新，还计划支持模式匹配、图分析和多模型数据访问等高级功能。

未来，随着 GQL 标准的完善和推广，预计将促进图数据库技术的进一步发展和应用。标准化的查询语言将简化图数据库的学习曲线，提高开发者的效率，并促进不同图数据库系统之间的互操作性，使得用户能够更容易地迁移和整合来自不同图数据库的数据。此外，标准化也有助于统一图数据处理的最佳实践，为图数据库的教育、研究和创新提供坚实的基础。总之，GQL 作为未来图数据库查询语言的标准，其发展将对整个图数据库领域产生深远的影响，推动图数据管理和分析技术向更高的成熟度迈进。

（二）图计算技术重要性

图计算技术是一种新兴的计算范式，它以图模型为基础，通过图算法对复杂关系数据进行分析 and 计算。近年来，图计算技术在科学计算和数据工程领域得到了广泛应用，并展现出巨大的潜力。

图计算目标是将计算从数值扩展到关系，传统科学计算主要基于数值计算，擅长处理结构化数据和求解确定性问题。然而，许多科学问题涉及复杂关系和不确定性。图计算技术能够有效地表达和处理复杂关系数据，并利用图算法进行分析和计算，为科学计算提供了新的思路和方法。

在众多领域，图计算不仅仅是一种技术选择，而是解决问题的必要手段。比如，在金融领域，图计算可以通过风险结构的特征数据描述，有效识别和防范欺诈行为，在经济上也可以通过图的流量特性，预测流动性风险等；在公安领域，图计算可以助力研究群体划分，分析人员活动等复杂交互作用，进而判断安全动向；在物流或交通领域中，图计算能够优化路径规划和资源分配。这些应用场景中，数据的关系性是核心问题的关键，而图计算正是解锁这一复杂度的钥匙。

图计算之所以必要，还在于它提供了一种能够与数据的自然结构相匹配的处理方式。在图中，数据以节点和边的形式存在，直接映射现实世界中的实体和它们之间的联系。这种一致性使得图计算在处理关系数据时更加高效和直观，无需进行复杂的转换或适配。此外，图计算框架和算法的设计允许直接在这些关系上执行复杂的查询和分析，从而大幅提升了处理速度和准确性。

随着技术的发展，图计算也在不断进化，以适应数据规模的增长和计算需求的提升。分布式图计算框架的出现，使得大规模图数据的处理成为可能，而图数据库的优化则提高了实时查询和分析的效率。这些进步不仅加强了图计算的必要性，也扩大了其应用范围，使其成为处理现代复杂数据问题的强大工具。

1.图计算技术特点

图计算技术的核心特点之一在于其能够实现大规模图数据的高效处理，主要得益于并行计算技术和分布式系统的结合。这种技术使得图计算不仅可以在单机环境中有效执行，更能在分布式环境中通过多节点协作，处理规模巨大、结构复杂的图数据。并行计算技术的应用极大地提高了图计算的性能，使得复杂的图算法能够在可接受的时间内完成计算任务，满足实时分析和决策的需求。

在并行计算模型中，图数据被分割成多个部分，分配到不同的计算节点上进行处理。这种数据分割策略关键在于如何有效减少节点间的通信开销和保证负载均衡，以提高整体计算效率。图计算框架通常采用基于顶点的编程模型，如 Pregel、BSP 或 Giraph，

其中计算以顶点为中心，每个顶点负责处理与其直接相连的边和邻居节点的信息。这种模型天然适合并行处理，因为各个顶点的计算可以独立进行，只在必要时进行节点间的信息交换。分布式图计算框架通常采用消息传递机制来实现节点间的数据交换和状态同步。在 BSP(Bulk Synchronous Parallel) 计算模式中，它将计算分成一系列的超步(SuperStep)的迭代(iteration)，每个计算步骤（或 SuperStep 超步）结束时，节点间会交换消息，以传递必要的状态信息或计算结果。这要求框架实现高效的消息传递系统，以支持大量细粒度消息的快速处理。同时，全局同步点（Barrier）用于确保所有节点在进入下一计算步骤前达到一致状态，保证计算的准确性。

图计算技术的另一个显著特点是其对算法的优化。在并行和分布式环境中，传统的图算法需要被重新设计或优化，以适应并行处理的需求和减少跨节点的数据传输。例如，通过算法分解和任务调度技术，可以将复杂的图算法拆分成多个小任务，分别在不同节点上并行执行，从而实现算法的加速。图计算技术的特点不仅体现在其对大规模图数据处理的能力上，更在于并行计算和分布式处理技术的深度整合，以及对图算法的优化和适应性改进。这些技术特点共同确保了图计算能够高效、可靠地应对现代数据分析中的复杂挑战，为从社交网络分析到生物信息学等多个领域的研究和应用提供了强大的支持。

2.图计算发展趋势

图计算的发展趋势正体现在多个维度，旨在解决现有技术的局限性，提升处理效率，拓展应用场景，以及加强用户体验。

(1)算法和框架的创新

随着图数据规模的不断增长，现有的图计算框架和算法正面临着新的挑战。为此，图计算通过新的框架和计算模型优化图计算是一个正在探索的方向，以更高效地处理大规模图数据。这包括开发能够更好地利用硬件资源的算法，如利用更先进的并行计算模型理论，GPU 和 FPGA 的并行计算能力，以及设计更为灵活和可扩展的图计算框架，以支持动态图数据的实时处理。

(2)图数据库与图计算的融合

图数据库的发展正在与图计算技术越来越紧密地结合。传统上，图数据库主要关注数据的存储和管理，而图计算则专注于对图数据的复杂分析。现在，越来越多的图数据库开始内置图计算能力，或者与图计算框架紧密集成，提供从存储到分析的一体化解决方案。这种融合使得用户可以更方便地对图数据进行高效处理和分析，无需在不同系统之间迁移数据。

(3)云计算和服务化

图计算服务正逐渐向云平台迁移，多家云服务提供商已经开始提供图计算作为服务（Graph-as-a-Service）。这不仅降低了用户部署和管理图计算环境的复杂性，也为用户提供了弹性的计算资源，能够根据需求动态调整。此外，云服务还促进了图计算技术的普及，使得即使是资源有限的用户也能够利用强大的图计算能力。

(4)科学计算跨领域应用的扩展

图计算的应用领域正在不断扩展，除了传统的社交网络分析、推荐系统和知识图谱之外，图计算也开始被应用于生物信息学、网络安全、智能交通系统等领域。随着算法和技术的进步，图计算在这些领域的应用将解锁更多的潜力，帮助解决更加复杂的问题。

(5)图神经网络（GNN）的兴起

图神经网络（GNN）作为一种结合图计算与深度学习的技术，正迅速成为图数据分析的热门领域。GNN 能够学习图中节点和边的深层次特征表示，为图数据的分类、预测和聚类等任务提供了强大的工具。随着 GNN 研究的深入和应用的拓展，预计未来图计算将更多地利用深度学习技术，以实现更加精准和高效的图数据分析。

(6)融合人工智能

人工智能技术与图计算技术的融合将是未来发展的重要方向。人工智能技术可以用于图数据的挖掘和分析，提升图计算的智能化水平。

（三）图数据中台的融合发展趋势与关键技术

图数据中台是一个集成了图数据库、图计算和图分析功能的数据处理平台，旨在为企业提供统一的图数据管理和分析解决方案平台，以支持复杂的数据关系分析、知识图谱构建、推荐系统等应用、人工智能平台以及扩展智能化能力。图数据中台的架构通常包括数据层、计算层、服务层和应用层。

- 数据层：负责存储原始数据和图数据，可以采用分布式原生图数据库来实现高效的图数据存储和查询。
- 计算层：提供图计算和图分析的能力，支持各种图算法的执行，如社区发现、路径分析、节点重要性计算等。
- 服务层：为上层应用提供统一的数据访问接口和服务，包括图数据查询、图计算任务提交等。
- 应用层：基于图数据中台提供的服务，开发具体的业务应用，如推荐系统、风险探测、知识图谱构建等。

1.图数据中台大模型、向量化、时序技术融合

知识图谱、大模型、向量化、时序技术是近年来人工智能领域发展的几项重要技术。它们的融合可以有效提升图数据中台的智能化水平、计算效率和处理能力。图数据中台的融合发展趋势体现在将大模型、向量化和时序技术与图数据处理相结合，以提升图数据的分析和应用能力。

知识图谱与大模型融合，是一个非常契合的结合，在技术层面互相补全了彼此的不足，融合技术在自然语言、推理、泛化、可解释性等方面都得到了补充。

(1)自然语言：知识图谱以结构化的形式存储知识，但缺乏语义理解能力；大模型具备语义理解能力，但缺乏知识支撑，容易形成幻象问题。融合两者可以弥补彼此的不足，提升知识的表达和理解能力。

(2)领域知识：大模型虽然通过大量数据训练获得了广泛的知识，但仍可能存在知识盲区，特别是在一些专业或细分领域。知识图谱提供了丰富的结构化知识，能够为大模型补充这些领域的具体和精确知识，从而弥补知识盲区，提高模型在特定领域的表现。

(3)提高可解释性：大模型的“黑箱”特性是其普遍受到的批评之一。通过与知识图谱融合，模型的决策过程可以依托于图谱中明确的实体关系和逻辑链，从而提供更加清晰的决策解释，增强模型的透明度和可解释性，这对于需要高度准确性和可信赖性的应用场景尤为重要。

(4)提高泛化能力：知识图谱可以为大模型提供普适性的知识和规则，帮助大模型学习更通用的模式，从而提升其泛化能力。

向量化技术融合，向量化技术包括节点嵌入和图嵌入，是图数据中台处理和分析图数据的关键技术之一。通过将图中的节点和整个图转化为低维稠密向量的形式，向量化技术使得图数据可以被应用于各种机器学习和深度学习模型中。这种技术的融合，不仅提高了图数据的可操作性，也为图数据的相似性度量、模式识别和聚类分析等任务提供了强大的支撑。

时序技术融合，时序技术的融合解决了图数据中台处理动态图数据的需求。动态图数据在多个领域中都非常常见，如社交网络的关系动态变化、金融市场的交易网络等。通过融合时序技术，图数据中台能够有效地处理和分析图结构随时间的变化，捕捉到图中的动态模式和趋势。这不仅增强了图数据中台的时效性和预测能力，也为用户提供了更为丰富和深入的洞察。

技术融合的实践意义，大模型、向量化和时序技术的融合，为图数据中台带来了前所未有的处理能力和应用可能性。通过这些技术的整合，图数据中台不仅能够处理规模更大、结构更复杂的图数据，还能够提供更为深入和精确的数据分析结果。在实践中，这意味着企业和研究者可以更好地利用图数据解决实际问题，无论是在推荐系统的个性

化推荐、金融风险管理的风​​险预测，还是在生物信息学的蛋白质相互作用网络分析等方面，都能够实现更高效和准确的结果。

2.图数据中台的实施步骤与挑战

图数据中台的实施通常包括以下几个步骤：

- (1) **需求分析**：明确企业的业务需求和数据特点，确定图数据中台的功能范围和技术选型。
- (2) **数据整合**：将分散在不同数据源的数据集成到图数据库中，构建统一的图数据模型。
- (3) **平台搭建**：根据架构设计，部署图数据库、图计算引擎等核心组件，搭建图数据中台的基础设施。
- (4) **功能开发**：开发图数据查询、图计算和图分析等功能，提供服务接口供上层应用调用。
- (5) **应用集成**：将图数据中台与业务系统集成，开发具体的图数据应用场景，如知识图谱应用、风险探测系统等。

在实施图数据中台的过程中，企业可能面临以下挑战：

- (1) **数据质量问题**：图数据的质量直接影响图分析的准确性，需要有效的数据清洗和融合方法保证数据质量。
- (2) **技术复杂性**：图数据中台涉及多种技术的集成，包括数据库技术、计算框架、大数据处理等，技术实现复杂。
- (3) **性能瓶颈**：随着图数据规模的增大，性能成为制约图数据中台应用的关键因素，需要不断优化存储和计算效率。
- (4) **安全与隐私**：图数据中台需要保证数据的安全性和用户隐私，遵循相关法律法规和标准。

图数据中台作为一种新兴的数据处理平台，通过融合大模型、向量化和时序技术，为企业提供了强大的图数据管理和分析能力。然而，在实施过程中也面临着数据质量、技术复杂性、性能瓶颈和安全隐私等挑战。企业需要根据自身的业务需求和技术基础，合理规划和实施图数据中台，以充分发挥其在图数据应用中的价值。

二、图数据库的发展趋势

随着数据规模的不断扩大和数据关联性的增强，图数据库在数据处理领域的重要性逐渐凸显。图数据库以其独特的图结构存储和查询方式，为数据处理带来了前所未有的灵活性和高效性。本章将探讨图数据库的发展趋势，以及它们如何塑造未来的数据管理。

(一) 市场增长与行业应用

随着各行各业对数据关联性需求的增加，图数据库的市场需求呈现出持续增长的趋势。据《2023 年全球图数据库市场发展空间分析报告》统计，在 2022 年全球图数据库市场规模达到了 116.1 亿元人民币（约合 17.7 亿美元），预计在接下来的预测期内，即 2023 年至 2028 年，该市场的年复合增长率（CAGR）将为 22.68%，到 2028 年市场规模将进一步扩大到 395.8 亿元人民币。2022 年中国图数据库的市场规模为 13.55 亿元人民币，占据了全球图数据库市场中相当大的份额。此外，中国的图数据库市场在未来几年内预计将继续保持稳健的增长势头，到 2028 年的预期市场规模为 65.3 亿元人民币。综上所述，全球图数据库市场正在迅速发展，而在中国这个市场已经取得了显著的进展。随着技术的发展和应用领域的扩展，预计图数据库的需求将会继续增加，从而推动整个市场的扩张。

金融、医疗、零售和电子商务等行业对图数据库技术的需求增加，进一步推动了市场的发展。不同行业对图数据库的需求有共性和差异。共性需求包括高性能、可扩展性、安全性和易用性。企业在选型时通常关注数据量、查询复杂性、实时性要求以及成本等因素。此外，行业特定需求也会影响选型，例如，金融领域更关注反欺诈能力，而医疗领域可能更关注数据隐私和合规性。

此外，随着数据规模的扩大，企业需要一种能够高效处理大规模数据的数据库技术。图数据库由于其高性能的查询处理能力，正逐渐成为处理大规模数据的重要工具。未来，随着数据处理需求的不断增长，图数据库的市场份额将进一步扩大。

(二) 技术成熟与迭代

图数据库目前为止应该经历了 4 个阶段，第一阶段是“非原生图数据库”阶段，这一阶段是随着关联分析、路径检索需求的不断出现，传统的关系型数据库为了满足海量数据关联分析性能要求，在关系型数据库之上增加了“逻辑图模型”而出现的支持图查询的“图数据库”，如 Oracle 中支持图查询与图分析。第二阶段应该是“原生图数据库”阶段，这一阶段是随着 NoSQL 数据库的兴起，以图模型作为数据模型的原生图数据库系统逐渐发展起来，如 Neo4J，gStore 等；第三阶段应该是“整合大数据处理平台的图数据库”，这一阶段随着 Hadoop、Spark 等大数据处理平台的兴起，图数据库系统开始与现有大数据处理平台开始进行深度整合，如 JanusGraph、GraphScope 等；第四阶段应该是“多模态的 HTAP 图数据库”，结合声音识别、图像识别等技术构建多模态图数据库，同时结合图计算引擎，实现“存算一体”的、支持事务分析混合处理的图数据库系统。

目前各个阶段都有一些产品，也都在各自不断发展，所以严格意义来说这几个阶段没有太明显的递进关系，更多的可能是不同发展方向。

此外，随着人工智能和机器学习技术的不断发展，图数据库将进一步集成这些技术，实现自动化数据处理和分析。通过机器学习和人工智能技术，图数据库将能够自动识别数据模式、预测数据趋势，为决策提供支持。

(三) 流批一体与分布式存储

为了更好地满足各行各业的需求，未来的图数据库将更加注重易用性和可靠性。未来图数据库技术将重点发展 HTAP 和流批一体技术，HTAP 就是事务分析混合型图数据库，既可以应用于事务型数据库场景，也可以应用于分析型数据库场景，而流批一体是指可以对数据进行批处理，即处理静态的、历史的数据集，也可以对数据进行流处理，即实时地处理一些数据流，实时产生结果，以满足更多场景需求。

然后是分布式技术图数据库技术，当前图数据的规模越来越大，已经达到千亿、万亿的处理需求，而且还会继续增长。学术界和工业界已经构建了不少分布式图数据管理系统，但随着大数据处理技术的不断深入发展，与大数据处理平台协同的分布式技术仍然是一个重要的发展趋势。

再者，图数据库在一些特定应用场景和运行环境下的技术也是一个重要趋势。例如动态图环境、低资源嵌入式环境、软硬件协同、时序数据等等，都面临不断增长的新需求，相应地对技术发展提出要求。

(四) 机遇与挑战

当前国内外图数据库市场均处于起步阶段，近些年一直在保持快速增长。首先，像金融、IT、电信等行业因为有大量数据和迫切的需求场景已经有成熟应用，产生了显著的经济和社会效益，也展示出了巨大的市场潜力；其次，整个社会对于图数据技术的认知也在逐渐扩展，因为基于实体和关系的图数据模型符合现实世界的抽象，相应的图分析方法也具有普适性，各行各业都可以去挖掘应用场景，未来想象空间很大；再者，大模型等相关技术发展产生的影响，也会进一步促进图数据库的普及和应用，大模型技术将自然语言处理能力大幅提升，对于非结构化数据转化为图数据有极大增强，两者在技术上具有互补性，可以相互促进。

当前图数据库确实机遇与挑战并存，因为处于起步阶段，众多厂商野蛮生长，机会很多。不过也面临多方面挑战：一是图数据行业相关标准还未健全，厂商之间存在多方面差异，导致图数据的交换困难，还未做到互联互通，而且多种图数据模型（RDF、属性图）与查询模型（SPARQL、Cypher、Gremlin、GQL 等）并存，难以统一应用；

二是图数据库的应用场景需要充分挖掘，需要找准图数据库的切入点，并且充分发挥图数据库在关联分析等方面的优势；三是将现有数据转化为图数据存在一定的技术挑战，工程量也较大。

在图数据库领域，国内外科研领域都取得很多优秀成果，市场上也都涌现出很多开源的、商业的优秀产品，属于齐头并进的形势。像北京大学自主研发的基于 RDF 的 gStore 图数据库系统，也赢得了很多国际同行的认可，在图查询优化、软硬件协同的图数据库查询加速等技术方面已经达到了国际领先水平。

同时，学术界和工业界关注点确实存在差异，学术界更关注核心、关键的单点技术的深挖和攻关，而工业界受市场驱动更关注图数据库产品的工程化和应用，需要考虑用户应用的方方面面，也更关注实际应用所产出的效益。不过就图数据库领域而言，工业界和学术界很早就进行联合，像蚂蚁、阿里、腾讯、华为等这些大厂都有自己的研究部门，而且和科研院所都建立了非常紧密的合作，或者是将高校的科研成果进行市场转化，在很多领域的成功应用中，都可以见到此类案例。如北京大学和华为建立了联合实验室，所取得的成果已经应用在华为数据通信部门的下一代路由器等嵌入式通讯设备上，大幅提升了核心技术指标，巩固了我国数据通信产品在国际上的技术竞争优势。因此，产学研用的联合已经被证明是非常有效的发展模式，对领域技术创新、成果落地以及合作各方都有促进作用。

总结：

随着数据规模的不断扩大和数据关联性的增强，图数据库的发展趋势愈发明显。市场增长与行业应用、技术成熟与迭代、查询计算一体化以及流批一体与分布式存储等趋势正塑造着未来的数据管理。在未来，图数据库将继续发挥其优势，为各行各业提供高效、灵活的数据处理服务。

三、图数据库的技术趋势

（一）图与大模型的结合

大模型（Large Model）是指运用深度学习技术构建的大型人工智能模型，通常具有大量的参数和复杂的网络结构。这些模型通过在海量数据上进行训练，能够学习和模拟高度复杂的数据模式和人类行为。它们通过在大规模文本数据上的训练，学会了理解和生成自然语言，可以进行问答、文本生成、翻译等多种任务。大模型的显著特点是它们的参数规模非常庞大，可以达到数十亿甚至数千亿级别。这使得它们在处理复杂问题时表现出色，能够捕捉到细微的数据规律和语言的微妙变化。

1.图与大模型结合的必要性

大模型在实际应用中也面临着诸多挑战,如训练数据的时效性问题、领域知识问题、幻觉问题、可解释性/溯源问题、推理问题等。

(1) 数据时效性问题: 大型语言模型通常在某一特定时间点完成训练,这意味着它们包含的信息是截至那个时间点的。因此,对于最新的事件、趋势或者研究成果,这些模型可能无法提供最新的信息。

(2) 领域知识问题: 一,数据集覆盖范围不够,大型语言模型通常是通过大规模的文本数据集进行训练的,这些数据集可能在某些专业领域的覆盖不够全面。例如,一些较为小众或尖端的科学领域、专业技术、行业特定知识可能在训练数据中的比例较低。二,难以应对专业知识的复杂性和动态性,一些领域(如医学、法律、高级工程学)涉及高度专业化和复杂的知识体系,且这些领域的知识在不断更新和发展。大型语言模型可能难以准确掌握和更新这些领域的深层次知识和最新发展。三,专业术语和概念的理解不深,专业领域往往有其特定的术语和概念。模型可能在理解和正确使用这些专业术语方面存在限制,尤其是在那些需要深入专业背景知识的场合。四,数据质量和准确性低,在一些专业领域,模型训练所依赖的数据可能存在质量问题,比如信息过时、来源不可靠、准确性有限等,这可能导致模型在这些领域的表现不佳。

(3) 幻觉问题: 在生成文本时,语言模型有时会创造出不存在或不准确的信息,这通常被称为“幻觉”。这可能是由于模型在训练过程中的数据限制或者理解上的局限性所致,导致“一本正经地胡说八道”。

(4) 可解释性/溯源问题: 一,缺乏明确引用,当模型生成文本时,它并不提供具体的来源或引用。这意味着用户无法直接知道信息的出处,无法核实其准确性和可靠性。这在需要依赖精确数据或权威信息源的情况下尤为重要。二,混合和重构信息,模型在处理和生成文本时,通常会综合多个来源的信息。这种混合过程使得单个生成的文本段落可能融合了来自不同来源的信息,而这些信息的原始出处在生成过程中丢失。三,无法区分事实与虚构,由于缺乏明确的引用和来源,模型生成的内容可能将事实、公认的信息与虚构或误解的信息混淆在一起。用户可能难以区分这些不同类型的信息。四,数据集来源的多样性,在训练过程中,模型使用了来自互联网的大量文本数据,这些数据来源广泛,包括新闻文章、书籍、网站内容等。但具体到每个回答或生成的文本,模型无法指明具体依赖了哪些数据源。五,对高度专业化内容的限制,在处理高度专业化的主题时,由于无法提供具体的学术或专业引用,模型的回答可能缺乏必要的准确性和深度。

(5) 推理问题: 大模型尽管在模式识别和关联建立方面表现出色,但它们在进行深层次、抽象或复杂逻辑推理时可能不够准确。一,对复杂推理链的处理有限,大型模型很擅长捕捉和模仿短期内的语言模式和逻辑关系,但在处理需要长期逻辑连贯性和多步骤

推理的任务时，效果可能不尽如人意。例如，它们可能在解决复杂的数学问题或进行深入的科学分析时遇到困难。二，抽象概念的理解不足，这些模型在理解抽象概念、原理和理论时存在限制。它们通过分析大量文本数据来“学习”这些概念，但可能无法真正理解或内化这些概念的深层含义和原理。三，逻辑演绎能力有限，尽管这些模型在关联和模式识别方面表现出色，但它们在逻辑演绎和推理方面的能力是有限的。它们可能无法有效地从已知事实推导出新的、正确的结论，特别是在信息不充分或需要高度逻辑推理的情况下。四，理解因果关系的挑战，大型模型在理解复杂的因果关系时可能遇到挑战。它们可能在辨别因果链中各个环节之间的逻辑联系方面不够准确，这会影响其在进行深入分析和预测时的效果。五，处理假设性和条件性语句的挑战，在处理包含假设性或条件性的语句时，模型可能无法准确捕捉其中的逻辑关系。例如，在处理“如果...那么...”类型的语句时，模型可能无法正确理解条件和结论之间的逻辑关系。

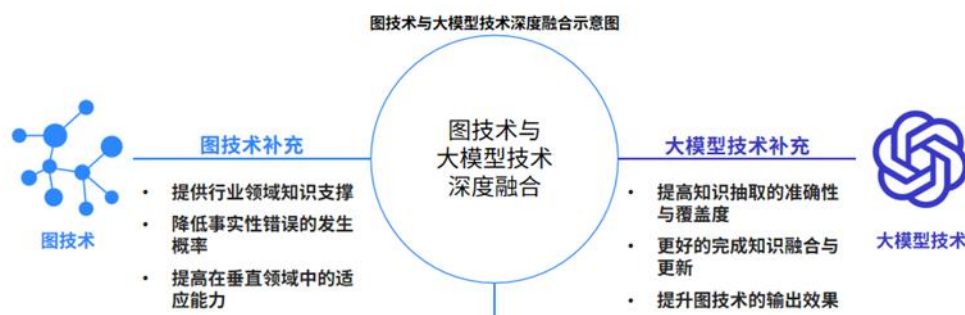
2.图与大模型结合的方法

本质上来讲，图和大模型都可以视为对知识的一种表达方式：大模型以参数形式表达知识，是对知识的一种隐式、非确定性表达；而图以实体和关系表达知识，是对知识的一种显式、确定性表达。因此，图和大模型在知识表达方面具有天然的互补性。

图擅长表示确定性的事实、时效性的知识/更新，可解释的推理等，大模型擅长语言理解与生成、对话式交互、跨语言知识获取、构建知识图谱、AI 编程等。通过图与大模型的共生方式，利用两者各自的优势相互赋能，弥补各自的不足，实现更高层次的认知能力。

在图模互补的共生系统中：大模型负责语言的理解与生成，实现对话式的交互，并负责跨语言的知识对齐和知识获取。具体到大模型对图的贡献上，则可以利用大模型在语义理解、内容生成等方面的技术优势，实现大模型对图构建至应用全生命周期各环节的增强，提升效率和质量。例如，利用大模型从文本中抽取实体和关系，构建或更新图；利用大模型根据图生成自然语言描述或问答对话，应用或传播图。

图则负责确定性的事实与知识，提供实时或及时更新的新鲜的知识，负责确定性的演绎推理和谓词逻辑，以及对错误知识的及时编辑与纠正等等。例如，由于图在知识标准化、可解释性、可溯源性、可控性、新鲜和及时更新等方面的优势，可以通过图来增强大模型从训练到应用的多个环节，提升大模型的应用效果和推理结果的可用性。例如，利用图作为额外的监督信号或输入特征，提升大模型在预训练或微调阶段的表现；利用图作为外部记忆或参考文献，在推理或生成阶段对大模型进行指导、纠正和知识溯源。



图与大模型深度融合示意图

（二）分布式图数据库计算存储优化

1.分布式图数据库计算存储优化的必要性

随着数字化转型的加速,人工智能和大数据技术的发展,对于能够灵活处理大规模、多样化数据的需求日益增长,图数据库因其高效的数据组织和处理能力,成为应对这一挑战的理想工具。图数据的规模也越来越大,对分布式图存储、图查询、图计算的优化要求越来越高。目前分布式图数据库主要面临以下几个问题:

超级顶点和热点访问的挑战。

超级顶点会存在大量的关系,对关系进行过滤和聚合会产生庞大的计算量,使得响应时间变长,从而影响业务,导致超时等。

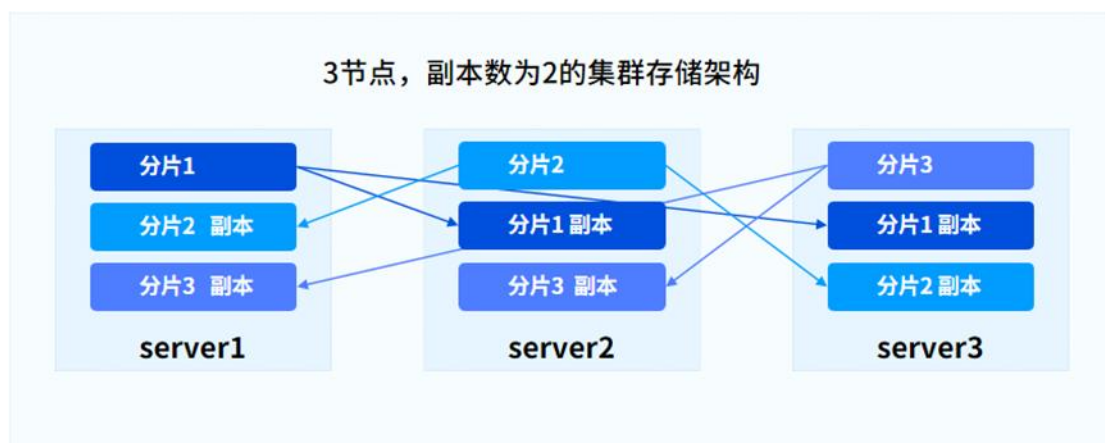
热点的高并发读写可能导致系统过载,难以在预期时间内处理请求,有时甚至会导致服务直接中断。

2.分布式图数据库计算存储优化的方法

分布式图数据库面对超级顶点和热点访问的挑战时,未来可能采取以下策略来解决这些问题:

- (1) **顶点切分**。超级顶点可以通过切分成多个较小的顶点来处理。这种方法可以减少单个顶点上的负载,使得关系的处理更加分散和高效。
- (2) **裁剪技术**。减少处理超级顶点时的计算负担。
- (3) **硬件加速**。使用 GPU 或者专门的硬件加速器进行计算,特别是在图的遍历和分析上。

通过这些策略,分布式图数据库可以更有效地处理超级顶点和热点访问带来的挑战,提高系统的可扩展性和稳定性,但值得注意的是,分布式数据库由于其存在跨节点的网络通信和额外的管理和同步的开销,特别是当查询涉及大量跨节点的数据交互时,其性能不一定比单机的图数据库高。分布式数据库解决的更多的是其扩展性、高并发和容错性。



分布式存储示意图

（三）图与时序数据的结合

1.图与时序数据的结合的必要性

时序数据是指随时间顺序记录的数据，它通常用于跟踪和记录在连续时间点或者时间段内发生的事件或测量值。这类数据特别适用于需要分析时间变化或趋势的场景，如股票市场数据、气象记录、物联网设备的传感器数据等。将图数据库与时序数据的结合可以提供强大的分析能力，不仅能帮助我们理解当前时刻的世界及其内在的复杂关联，还能揭示这些关系是如何随时间演化而来的。这种结合提供了一种独特的视角来分析和理解数据，特别是在那些动态变化和发展对决策过程至关重要的场景中。

2.图与时序数据的结合的方法

- 时序数据作为图属性，在图数据库中，可以将时序数据作为顶点或边的属性。例如，一个社交网络中的用户顶点可以包含用户活动的时序记录，或者一条交易边可以包含交易金额随时间的变化。
- 时序图分析，通过对图中的时序数据进行分析，可以识别时间相关的模式和趋势，比如发现在特定时间周期内重复出现的模式。
- 时间切片分析，结合图数据与时序数据分析各个时间切片上的图数据，这种强大的分析手段能够揭示随时间演变的复杂关系和模式。这种方法特别适合于那些动态数据和网络结构对决策至关重要的领域，如社交网络分析、金融市场监控、供应链管理等。
- 时间序列预测与图关系，结合时间序列分析方法和图的关系数据，可以进行更复杂的预测，比如预测一个节点的未来状态或者一种趋势如何在图中传播。
- 事件驱动的图更新，时序数据可以用于触发图的更新，例如在监控系统中，一系列时间点上的事件可能会导致图中的节点或边的创建、更新或删除。

通过将图数据库与时序数据结合，可以极大地增强数据的表达力和分析能力，尤其是在需要理解复杂系统内部随时间演变的动态关系时。这种结合在金融分析、网络安全、物联网、社交网络分析等领域有着广泛的应用。



社群演进时序分析

四、图数据库测试标准

图数据库是未来数据管理的发展趋势之一，通过图数据库测试来评估图数据库能力至关重要。现如今，图数据库测试方案多种多样，但存在各种弊端。所以需要一个权威、真实、公平的图基准测试。关联数据基准委员会（LDBC, Linked Data Benchmark Council）是由 Oracle、Intel、蚂蚁集团等软硬件巨头，Neo4j、TigerGraph、创邻科技、海致星图、Ultipa 等国内外主流图数据库厂商，伦敦大学、爱丁堡大学、希腊研究与技术基金会等国际知名学术组织等组成的非营利机构，是全球图数据库领域唯一的第三方、非营利、权威测试标准制定与发布机构，在行业内有着很高的影响力。LDBC 提供了公平、公正、公开的图基准测试标准，LDBC 的图基准测试以图模式查询中常见的瓶颈点（Chock Points）出发，在查询计划优化、查询执行引擎、语言表达能力等多个方面，从聚合、连接方式、数据访问方式、表达式计算、子查询优化、并行化、图特征计算、更新操作等多项维度出发，基于相同的标准、相同的数据集、相同的测试工具，全面考察图数据库的性能、高可用、并发等能力。

（一）LDBC 对图模式查询的瓶颈点设计

LDBC 对图模式查询的瓶颈点设计主要覆盖 5 个方面：查询优化器、执行引擎、存储引擎、查询语言表达能力、数据更新。具体的分类包括：

1. 聚合：包括排序、并行排序、TopK 下推等，比如追踪资金流向业务中，找出最近 10 次转账；
2. 连接方式：包括深度优先遍历 (Depth First Search)、广度优先遍历 (Breadth First Search) 的图遍历方式，比如追踪资金流向业务中，找出 5 跳内的转账链路；
3. 数据访问方式：包括随机查询的能力，比如根据 id 查找顶点，根据顶点找到邻居，获取邻居的属性信息；
4. 表达式计算：包括公共子表达式复用的能力，比如统计杭州一天的交易总量；
5. 子查询优化：包括一个查询内相同或相反的计算，合并到一个子查询的能力、不同查询间，复用子查询结果的能力、一个查询内复用子查询的能力；
6. 并行化：包括在大量查询的情况下，使用缓存子查询的能力；
7. 图特征计算：包括基于图遍历重用之前找过的路径的优化能力、基于查询两点间路径的优化能力、基于查询两点间不带权最短路径的优化能力、多个图查询组合查询的能力等；
8. 更新操作：包括对顶点、边以及属性的增删改查操作。

(二) LDBC 图基准测试

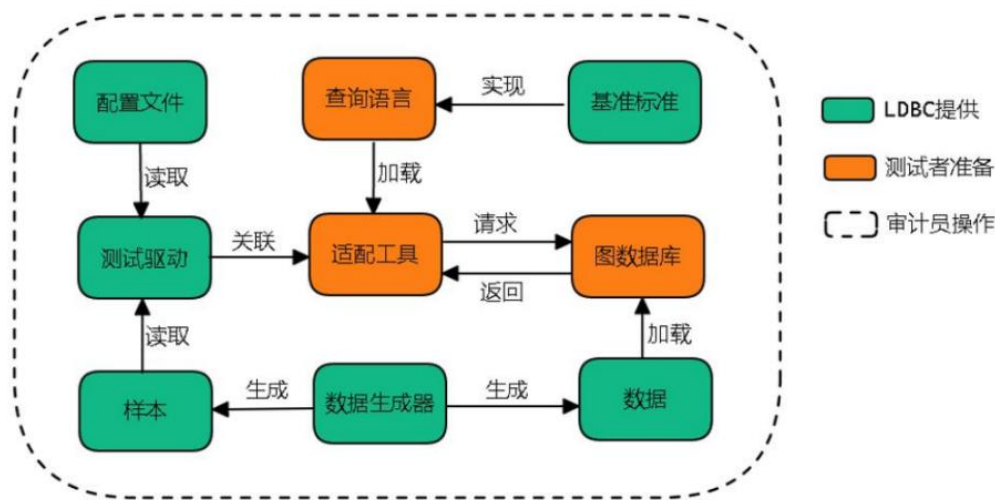
LDBC 图基准测试基于应用层级测试、面向瓶颈点的原则进行设计，目前有 LDBC SNB、LDBC FinBench、LDBC SPB、LDBC Graphalytics 面向不同的图数据库和图计算系统。其中：

1. LDBC SNB 是面向属性图图数据库的基准测试，总结了社交场景下的数据特征和负载特征，对图数据库系统展开测试。LDBC SNB 定义了两个 Workload (测试负载)，Interactive Workload 和 BI Workload。其中 Interactive Workload 定义了 TP 场景的测试负载，BI Workload 定义了 AP 场景的测试负载。
2. LDBC FinBench 总结了金融场景下的数据特征和负载特征，面向金融场景对图数据库系统展开测试。LDBC FinBench 由蚂蚁集团发起，由 Intel、TigerGraph、创邻科技、Ultipa、Nebula 等国内外主流图数据库厂商共建发布。LDBC FinBench 目前仅定义了 Transaction Workload，面向 TP 场景。
3. LDBC SPB 是面向 RDF 图数据库的基准测试，基于 BBC 出版场景定义了增删查混合的负载，对 RDF 图数据库展开测试。

4. LDBC Graphalytics 是面向图计算系统的基准测试，从标准图算法在不同的数据规模上对图计算系统的展开性能测试。

(三) LDBC 图基准测试方式

LDBC 图基准测试提供了标准的测试工具，相当于黑盒测试工具，测试方无需自己实现测试工具的逻辑，确保测试方式一致。如下图 LDBC 图基准测试流程所示，黄色是测试者需要准备的内容：即查询语言，适配工具和测试的图数据库产品，LDBC 的官方会提供数据生成器与测试驱动的工具，包含了数据生成、样本生成、正确性验证、性能测试的功能，在整个测试过程中，LDBC 审计员会全程审计测试过程，严格记录每一个过程的日志和结果并公开，是一套非常完善的基准测试流程。



图：LDBC 图基准测试流程

9、文档数据库

一、文档型数据库的兴起

文档型数据库，作为 NoSQL 技术最重要的部分，在数字化、AI 智能化对于更高灵活性和扩展性的数据管理领域需求中，越来越重要。在处理大规模、多样化、非结构化或半结构化数据方面，它们提供了一种高效且灵活的解决方案。文档型数据库的兴起，标志着从严格的表格型数据结构向更自由的数据组织模式的转变。

(一)技术背景

随着互联网技术的飞速发展和数字化数据的爆炸式增长，传统的关系型数据库（RDBMS）在处理大量、多样化的数据方面逐渐显现出局限性。与此同时，开发者和企业对于数据库系统的要求也越发多元化，他们寻求一种可以快速适应业务变化、支持灵活数据模型的数据库解决方案。

（二）文档型数据库的特点

1. **灵活性**：文档型数据库允许数据以 JSON（JavaScript Object Notation）、BSON（Binary JSON）、XML 或类似的格式存储，这些格式支持层次化的、嵌套的数据结构，为数据的快速开发和迭代提供了便利。
2. **易用性**：文档型数据库通常提供与其数据格式紧密结合的查询语言，使得数据的读取和写入更为直观和高效。
3. **扩展性**：它们天生支持分布式架构，可以方便地扩展到多个服务器，以处理更大规模的数据集和更高的并发需求。

（三）市场响应

得益于其数据结构的灵活性，市场对文档型数据库的接受程度很高，特别是在需要快速开发和部署的创新型企业和互联网行业中。这些数据库被广泛应用于内容管理、电子商务、移动应用开发等领域，充分体现了其在现代应用中的重要性。

我们可以将文档型数据库比作一本内容丰富、结构灵活的杂志，与传统的按章节严格划分的书籍（关系型数据库）相比，杂志及新媒体（文档型数据库）可以容纳各种格式和类型的内容，更适合快速浏览和更新。

二、文档型数据库的历史背景和初期发展

（一）早期探索

文档型数据库的历史可以追溯到 20 世纪 80 年代末至 90 年代初，这是计算机科学历史上的一个转折点，当时关系型数据库（RDBMS）在商业和学术界占主导地位。然而，随着互联网的兴起和数据量的激增，关系型数据库的限制开始显现，特别是在处理大量的非结构化数据方面。这促使了对更灵活、更可扩展的数据库解决方案的探索，最终导致了文档型数据库的诞生。

（二）技术革新

文档型数据库的早期开发集中在提供一种不同于传统表格型数据结构的方法。这种新型数据库使用文档（如 JSON、XML、BSON 格式）来存储数据，每个文档都可以

有自己独特的结构。这种灵活的结构使得文档型数据库能够轻松处理多种类型和格式的数据，特别适合快速迭代和开发，这在当时的数据库技术中是一个重大的创新。

（三）关键发展阶段

1. **Lotus Notes**：在 90 年代初，Lotus Notes 是最早采用文档型数据模型的系统之一。它不仅作为一个电子邮件客户端，还提供了文档数据库的功能，允许用户以文档的形式存储和共享信息。
2. **开源运动**：21 世纪初，随着开源运动的兴起，文档型数据库技术开始迅速发展，CouchDB 和 MongoDB 是全球著名的例子。而在中国主要有 SequoiaDB 巨杉数据库于 2014 年开源首个版本，并且已经在超过 100 家金融银行总部得到广泛使用，多用于支持银行影像数据平台、非结构化数据平台等企业内容管理系统。它们都提供了更高级的查询功能、更好的性能和更强的扩展性。

（四）挑战与解决方案

早期的文档型数据库面临着多种挑战，包括数据一致性、查询效率和数据安全等问题。为了解决这些问题，开发者和企业在后续的版本中引入了更复杂的索引机制、更强大的查询语言和更严格的安全措施。

（五）市场接受度

尽管初期存在挑战，文档型数据库因其在应对非结构化数据处理方面的优势而迅速受到市场的欢迎。它们特别受到内容管理系统、电子商务平台和大数据应用的青睐，因为这些领域经常需要处理各种格式和类型的数据。

如果将关系型数据库比作传统图书馆中严格分类的书籍，那么文档型数据库就像个人书架，可以自由地存放各种类型的书籍和杂志，甚至是手写笔记。这种自由和灵活性使得个人书架（文档型数据库）非常适合迅速变化和多样化的现代信息需求。

三、文档型数据库在中国的发展

随着中国数字经济的快速发展和技术创新的推进，文档型数据库在中国市场也经历了显著的成长和变化。

（一）市场环境和驱动力

1. **数字化转型**：随着中国企业数字化转型的加速，对于能够有效管理和利用大数据的需求日益增长。文档型数据库以其高效的数据处理能力和灵活的数据模型，成为众多企业数字化战略的关键组成部分。

2. 技术创新：中国本土企业在云计算、大数据、人工智能等领域的持续投入和创新，为文档型数据库的应用提供了强大的技术支持和广阔的应用场景。

（二）发展历程

1. 初期引入：早期，中国市场上的文档型数据库多为国外产品，如 MongoDB 和 CouchDB 等。
2. 本土化发展：随着市场需求的增长，国内企业如 SequoiaDB 开始开发本土化的文档型数据库产品。
3. 与中国关系型数据库发展不同的是，关系型数据库全球发展已经在全球商业化发展近 50 年，而中国商业化产品发展仅 20 年。而在文档型数据库方面，中国和全球几乎同步发展，MongoDB 成立于 2008 年，与成立于 2012 年的巨杉文档型数据库仅差 4 年，而他们的商业化产品均在 2013 年正式发布完全同期发展。

（三）关键应用领域

1. 内容管理系统（CMS）：文档型数据库在内容管理系统中被用来存储各种格式的内容，如文本、图片和视频等。如 MongoDB 及巨杉文档型数据库更专门推出 GridFS 及 LOB 对象管理能力，承载海量的非结构化对象数据。
2. 移动互联网：随着移动互联网的蓬勃发展，许多移动应用选择文档型数据库来存储和处理大量的用户数据和日志信息。

四、文档型数据库的技术演进与创新

随着时间的推移，文档型数据库在技术上逐渐成熟，它们开始提供更加强大且多样的功能，以应对不断增长的数据管理需求。

（一）关键技术进步

1. 索引与查询优化：为了提高查询效率，文档型数据库引入了更加高效的索引机制。例如，系统采用了 B 树索引，这大大加快了查询速度，特别是在处理大量数据时。同时，这些数据库也增强了它们的查询语言，使得复杂查询、数据聚合和全文搜索变得更加容易和高效。
2. 数据模型的灵活性：文档型数据库的一个关键优势是它们的数据模型非常灵活。这意味着文档的结构可以随需要改变，不需要像关系型数据库那样预先定义固定的数据模式。这种灵活性使得文档型数据库非常适合快速开发和持续迭代的项目。

3. 分布式架构与扩展性：随着云计算的兴起和大数据时代的到来，文档型数据库开始采用分布式架构，以提供更好的扩展性和可靠性。这些数据库能够跨多个服务器和地理位置分布数据，从而提高数据处理能力和容灾能力。

（二）创新案例

1. MongoDB 的 Aggregation Pipeline：MongoDB 引入了 Aggregation Pipeline 功能，它允许开发者在数据库层面进行复杂的数据处理和转换操作，这在提升了数据处理效率的同时，也降低了应用服务器的负担。

2. 巨杉文档型数据库的精细化容灾：针对跨数据中心部署的业务场景，面向海量非结构化数据，通过 Location 亲和性配置策略提升系统故障切换管理的自动化能力，同时通过限流、数据域复制粒度，在数据可靠性及数据中心带宽控制方面提供可调节的平衡能力。

3. CouchDB 的复制功能：CouchDB 提供了强大的数据复制功能，使得数据可以在不同的服务器和设备之间轻松同步，这对于构建高可用和容灾强的系统至关重要。

（三）挑战与应对

尽管文档型数据库在技术上取得了显著进步，但它们仍面临着数据一致性和事务处理方面的挑战。为了应对这些问题，一些文档型数据库如：MongoDB、巨杉文档型数据库开始提供更加严格的事务控制机制，比如多版本并发控制（MVCC）和原子操作。

文档型数据库的技术演进，就如同一艘正在升级的船只。最初，这艘船只适合近海航行（简单的数据管理任务），但随着技术的发展，它被装上了更强大的引擎（索引优化）、更灵活的船舱设计（数据模型灵活性）和更先进的导航系统（分布式架构），使其能够承载更重的货物（大数据处理）并远航到更远的海域（复杂的应用场景）。

五、文档型数据库的行业应用与案例分析

文档型数据库因其灵活性、易用性和强大的数据处理能力，在多个行业中找到了广泛的应用。以下是几个关键行业的应用案例分析：

1.内容管理系统（CMS）：

- 应用场景：文档型数据库在内容管理系统中被用来存储各种格式的内容，如文本、图片和视频等。
- 案例分析：一家新闻门户网站使用文档型数据库来管理其新闻文章和多媒体内容。文档型数据库的灵活性允许快速调整内容结构以适应不同类型的报道和特色栏目。

2.物联网（IoT）：

- 应用场景：在物联网领域，文档型数据库用于处理来自传感器和设备的大量实时数据。
- 案例分析：一个智能家居系统使用文档型数据库来收集和分析来自各种智能设备（如温度传感器、安全摄像头、照明系统）的数据。数据库的高效性能确保了即时的数据处理和响应，提高了整个系统的智能化和用户体验。

3.移动应用开发：

- 应用场景：在移动应用开发中，文档型数据库因其轻量级和高性能而受到青睐，特别是在需要处理大量用户生成内容和实时数据的应用中。
- 案例分析：一款社交媒体应用利用文档型数据库来存储用户的帖子、评论和个人资料。文档型数据库的高效读写能力保证了应用即使在高用户负载下也能保持流畅的性能。

4.电子商务：

- 应用场景：在电子商务平台中，文档型数据库用于处理大量的产品信息、用户数据和交易记录。这些数据通常是非结构化或半结构化的，且需要频繁更新。
- 案例分析：例如，一个大型在线零售商使用文档型数据库来存储其庞大的产品目录。每个产品作为一个独立的文档，含有价格、描述、用户评价等信息。这种结构使得添加新产品、更新现有产品信息或实现个性化的用户推荐变得更加简单和高效。

5.大数据分析：

- 应用场景：文档型数据库在大数据分析中被用于快速处理和分析来自各种源的数据。
- 案例分析：一家营销分析公司使用文档型数据库来聚合和分析来自社交媒体、网站流量和客户数据库的数据。利用文档型数据库灵活地聚合和查询功能，公司能够快速生成定制化的市场趋势报告和消费者洞察。

文档型数据库在这些行业中的应用展示了它们在处理多样化、动态变化的数据集方面的独特优势。通过灵活的数据模型和高效的数据处理能力，文档型数据库为这些行业提供了强大的后端支持，帮助企业更好地理解数据，提高操作效率，创新服务和产品。他就像是一位多才多艺的艺术家：在内容管理系统中，它是一个创意无限的编辑；在物联网领域，它则是一个精通多种语言的翻译官；在移动应用开发中，它是一位快速反应的设计师；在电子商务领域，它是一位细心的商品陈列师；而在大数据分析中，它化身为一位敏锐的市场分析师。在各个领域，文档型数据库都能展现其独特的才华和价值。

六、文档型数据库的当前挑战与未来趋势

文档型数据库，尽管在许多方面取得了显著的成功，仍面临着一些挑战，同时也正向着新的发展趋势迈进。

（一）当前挑战

1. **数据一致性**：在分布式系统中保持数据一致性一直是文档型数据库的一个挑战。尤其是在多节点环境下，如何有效地处理数据的同步和冲突解决是关键问题。
2. **安全性和隐私**：随着数据泄露和网络攻击事件的增多，如何保护存储在文档型数据库中的敏感数据成为一个重要议题。实施强大的安全措施和符合法规的数据处理机制是文档型数据库必须面对的挑战。
3. **高级查询性能**：虽然文档型数据库在处理复杂查询方面已经取得了进步，但在某些情况下，它们的查询性能仍然无法与传统的关系型数据库相媲美，特别是在涉及大数据集的高级分析和报告生成方面。

（二）未来趋势

1. **人工智能和机器学习集成**：预计文档型数据库将更深入地融合人工智能（AI）和机器学习技术。文档型数据库也将内嵌 Vector 向量检索能力，可以成为 RAG 的基础数据底座，为企业私有化 GenAI 服务提供安全的数据边界及处理能力，以应用于自动优化处理、智能管理数据分布和自动化安全监控等广泛的业务领域。
2. **云计算和边缘计算的融合**：随着云计算的普及，文档型数据库越来越多地被部署在云环境中。同时，边缘计算的兴起预示着数据库技术将更多地被集成到离用户更近的位置，以减少延迟和提高数据处理速度。
3. **多模型数据库的发展**：预计将有更多的文档型数据库向多模型方向发展，即在单一的数据库系统中支持多种数据模型（如文档、键值、图等），为开发者提供更多的灵活性和选择。

10、时序数据库

一、时序数据库演进史

时序数据的主要产生来源之一是设备与传感器，具有监测点多、采样频率高、存储数据量大等多类不同于其他数据类型的特性，从而导致数据库在实现高通量写入、存储成本、实时查询等多个维度存在管理难点。针对这些特性与难点，专门针对时序数据管理构建的时序数据库也在逐步成熟。

（一）时序数据管理的早期方案

1.最早的“时序数据库”：RRDtool

时序数据库的起源可以追溯到 20 世纪 70 年代，随着工业控制和 SCADA 系统的兴起，大量的时序数据产生，于是需要一套完整的存储与处理方案。而 1999 年出现的 RRDtool（Round Robin Database Tool）最早提出了专门面向时序数据存储、处理的方法。RRDtool 命名中提到的 Round Robin 其实是一种存储数据的方式，使用固定大小的存储空间，并有一个指针指向数据库中最新数据的位置。

如果将这个固定的存储空间想象为一个类似时钟的圆盘，上面有很多代表数据存储位置的刻度，指针就可以看作像时针/分针一样，是一条从圆盘中心指向刻度的直线。就像时钟可以不停转动一样，RRDtool 中的指针也可以一直移动，在存储空间足够的情况下，不存在无法存储新数据的问题；而因为数据存储空间是固定的，当所有的空间都存满了数据，就会覆盖最老的数据。

2.RRDtool 的创新与不足

从这样的设计就可以看出，RRDtool 的存储方式具有明确的时间属性，所以它适合存储时序数据，它绘制出的图类型也非常契合时序数据的属性，即以时间为横轴，以数值为纵轴的折线演变，强化了时序数据“以时间为第一概念”的数据特性。同时，因为 RRDtool 的存储空间大小是已经被定义好的，当空间存储满后，它将从指针的开头开始重新存储，数据集不会增大，所以存储空间大小不需维护。

然而，这类数据库也能够很明显地看到它的问题。首先，因为存储空间的手动设定，RRDtool 的存储能力难以扩展，在数据量逐渐增多的情况下，很难覆盖历史数据，实现时序数据的“应存尽存”。其次，RRDtool 的数据读取功能较弱，缺乏针对时间维度的查询优化，处理的数据模型也较为单一，通常是内嵌在监控系统中。最后，RRDtool 仅支持单机模式，未覆盖分布式管理的需求。

（二）时序数据库发展期：针对特征，优化性能

1.从 OpenTSDB 到 InfluxDB

随着大数据的发展，时序数据爆发式增长，已有方案逐渐不能满足需求，在 2010 年之后，出现了第二类产品，首先是以 OpenTSDB 为代表的基于分布式存储的时序数据库。这类时序数据库在继承通用分布式存储的基础上，扩展了时序数据的语义，并针对时序数据的查询、处理等进行了优化。如 OpenTSDB 底层依赖于 HBase 集群存储，根据时序的特征对数据进行压缩，节省存储空间；对时序数据的常用查询进行封装，提供数据聚合、过滤等操作。

而 OpenTSDB 存在的部署复杂和维护成本高等问题,促进了低成本的垂直型时序数据库的诞生,也就是以 InfluxDB 为首的时序数据库。这类时序数据库的目标场景是互联网的服务监控、运维等场景,拥有更灵活的数据模型,以标签模型对监控项进行管理。相对于 OpenTSDB 需要配置 Java 环境和 HBase 环境,InfluxDB 没有依赖,大大减少了开发与运维成本,易于部署和维护。

同时,InfluxDB 针对时序数据特性进行了存储引擎、查询引擎的重新设计,使得读写性能、易用性相比 OpenTSDB 获得了明显提升,如 InfluxDB 的采用类似 LSM Tree 的 TSM Tree 存储结构,引入了 series-key 的概念,根据时间特征对数据实现了很好的分类,减少冗余存储,提高数据压缩率。

2.时序数据库下一步发展挑战

上面提到的以 InfluxDB 为代表的时序数据库产品,在监控场景时序数据的多个管理痛点,如读写性能、存储成本、特性查询、部署运维等多个方向均实现了性能突破,并结合云服务、微服务等其他发展趋势,正在进一步拓宽其集成链路 with 适用场景。

然而,InfluxDB 主要面向云端服务监控,并且重点面向近期几个月的数据。但在时序数据大规模产生的工业物联网场景中,数据往往需要管理数年,甚至数十年之久,且需要分布式部署。InfluxDB 目前仅单机版开源,难以管理如此大的数据量,且随着存储时长的增加,查询性能会大幅下降。

在工业场景中,时序数据的管理还具有一些行业特点,数据大多是从端侧设备产生出来的,这些数据首先会服务于工厂的应用管理,所以它们会首先传到工厂内部的边缘网关,再传输到中心侧的数据库去支持监控、告警等服务。边缘侧网络资源、环境配置往往有限,此类时序数据库产品部署的带宽成本较高,对于数据库多终端之间同步传输的运维压力也较大。

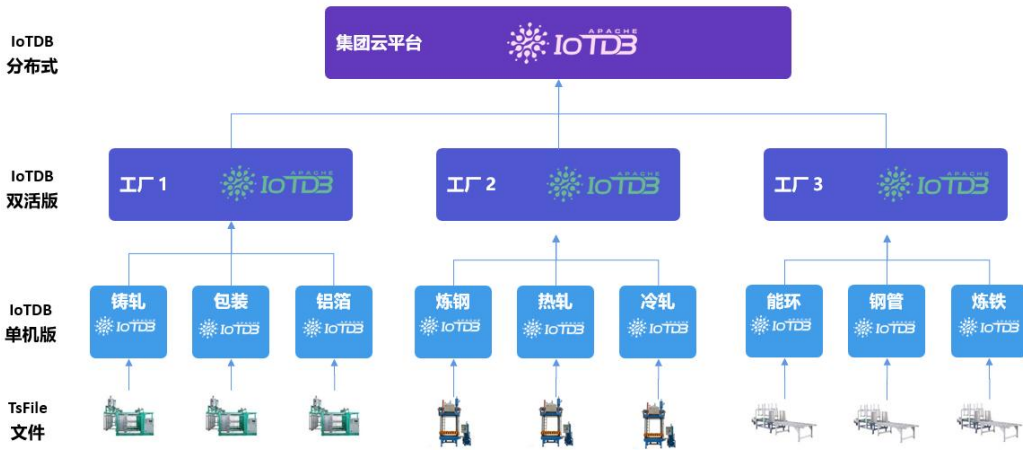
(三) 时序数据库的未来方向:实现工业物联网场景的“端-边-云”协同

1.工业物联网场景的端边云协同解决方案

面向工业物联网时代,国产自研的时序数据库产品也开始发展、成熟,试图打破垄断。而此类代表为技术发源于清华大学,历经超 10 年发展的物联网原生时序数据库 IoTDB。IoTDB 自研完整的存储引擎、查询引擎、计算引擎,并支持权限管理、集群管理、系统监控、可视化呈现等多项功能,具有多协议兼容、超高压压缩比、高通量读写、工业级稳定、极简运维等特点。IoTDB 可协助企业构建一体化时序数据收集、存储、管理与分析解决方案,其核心技术成功获得了相关发明专利 40 余项,并在 ICDE、

SIGMOD、VLDB 等数据库顶级会议上发表论文多篇。以 IoTDB 为代表的新的国产时序数据库旨在将 OT 和 IT 结合起来，提供完整的工业物联网时序数据解决方案，实现“端-边-云协同”，即需要时序数据库在端侧、边侧、云侧等不同资源下都能够适配，并且运行良好，可以进行数据管理和分析，同时在数据流转的过程中支持稳定、高效、低运维成本的数据传输方案。这一类时序数据库不但可以实现千万级数据写入、高压缩比数据存储、集群部署高可用等性能，同时通过独特的时序数据存储文件格式，实现了数据一次处理，即可端边云共用的新形态。

此类数据库在发送端可将数据持久化为文件格式，利用边缘端的计算能力，将数据进行深度编码压缩。当文件在端侧形成时，即可将其直接传输到接收端，而无需将原始数据传输。在接收端，此类数据库可将收到的文件直接加载至系统中，无需将数据解码和重新写入。同时，该类数据库通过“端-边-云”数据模式的自动识别，实现了低流量数据同步方案。因此可使底层数据文件格式贯穿端侧、边侧、云侧，支持可插拔的文件级同步，充分利用边缘计算能力，缓解云侧计算压力，有效节省网络流量、云端计算资源与运维成本。



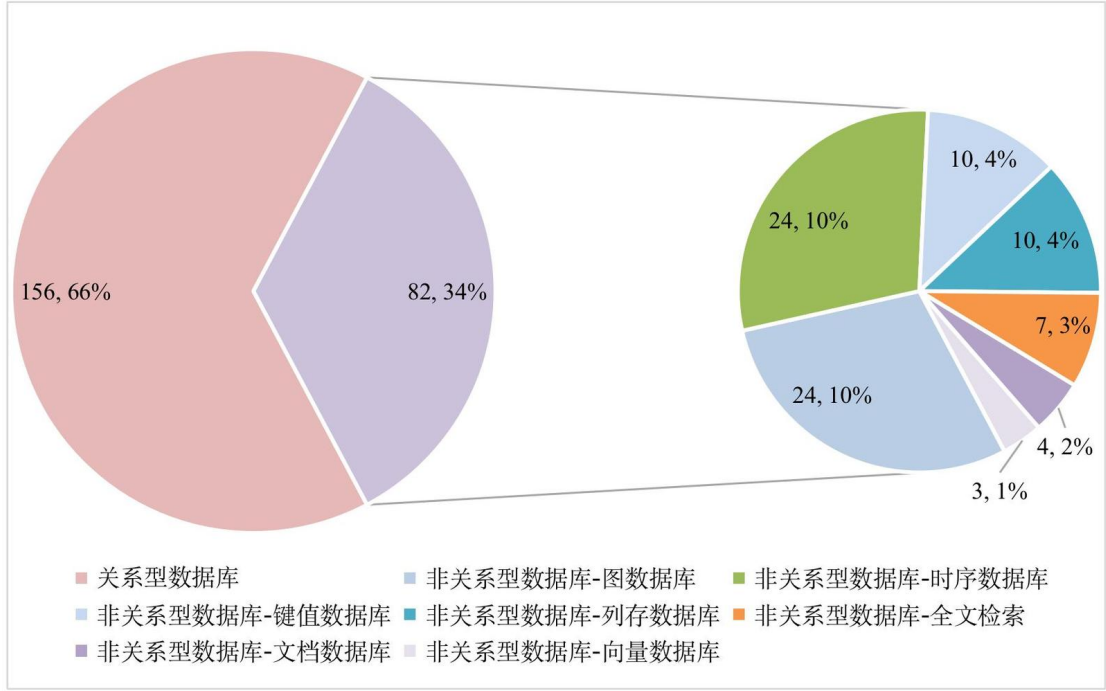
2.时序数据库端边云场景应用方向

随着工业 4.0 和数字化时代的到来，结合国家产业升级转型战略的实施，新一代的时序数据库产品开始深入地帮助工业企业实现深度数字化转型。以 IoTDB 为例，目前其广泛应用于能源电力、石油化工、钢铁冶炼、航空航天、轨道交通、智能工厂、车联网等国民经济核心产业，支持千万级设备写入、每日过亿条新增数据、TB 级历史数据存储需求，并实现如电力行业中“电厂侧-省域侧-中心侧”、制造行业中“传感器侧

-厂站侧-基地侧-集团侧”的“端-边-云”数据同步，构建单平台全生命周期与跨平台端边云协同的时序数据解决方案。

二、时序数据库发展趋势

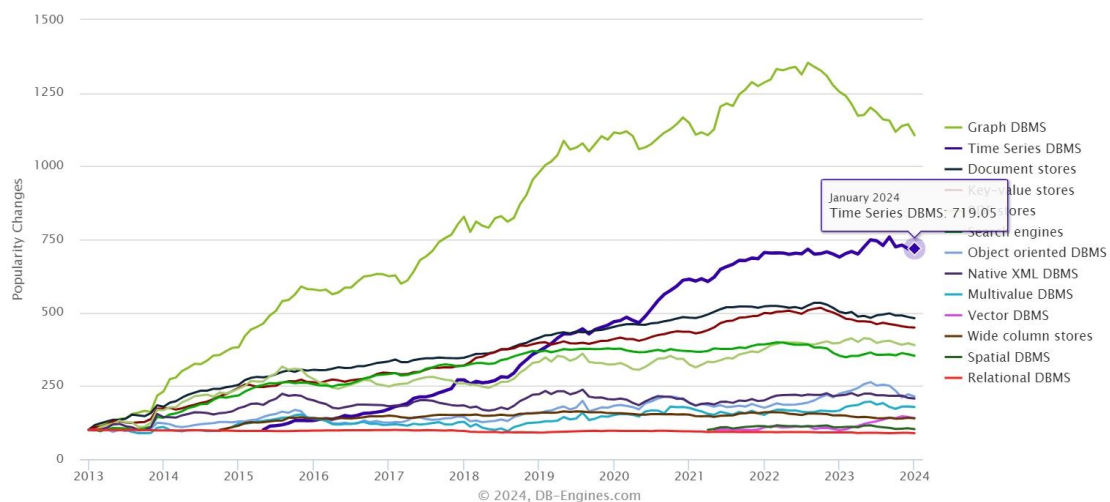
随着物联网、工业互联网、数字能源技术的快速发展，各类应用产生的时序数据量呈爆炸式增长。在这种背景下，时序数据库（Time Series Database，TSDB）迅速崭露头角，成为应对时序数据挑战的关键技术之一。据中国信通院发布的《数据库发展研究报告（2023 年）》显示，截至 2023 年 6 月，我国共有数据库产品 238 款，其中关系型数据库 156 款，非关系型数据库 82 款。在非关系型数据库中，时序数据库以累计 24 个产品位居第一，与图数据库并列，表明时序数据库在我国数据库市场中的重要地位及强劲的发展势头。



来源：CCSA TC601，2023 年 6 月

图 1：中国数据库产品类型分布

此外，根据 DB-Engines 的数据库发展趋势公开数据，时序数据库的流行度持续上升。从 2013 年到目前编写本文的时间（2024 年 1 月），时序数据库流行度一路高歌猛进，从 2020 年开始跃居第二位，仅次于图数据库。



来源：DB-Engines，2024 年 1 月

图 2：DB-Engines 近 10 年各模型数据库流行度趋势图

从时序数据库的技术特点来看，时序数据库的一大主要特点是：数据带有时间属性，具有高并发、快速写入、可扩展性等特点。以时序数据库典型应用场景 IoT 设备来看，其服务的设备数量多，生成数据快，数据量大。IDC 曾在《全球物联网设备数据报告》中预测，2025 年全球范围内将有 416 亿台物联网连接设备，每年产生的数据量将达到 79.4 ZB。因此，时序数据库的一个首要的技术方向即具有高并发的数据写入能力，通常需要每秒插入的数据量在 1 万以上，甚至可以达到每秒百万级别。另外一个技术方向是随着数据量的爆炸式增长，要求 TSDB 具有较强的扩展能力，支持数据的水平扩展和负载均衡部署。时序数据的另一个特点是：数据是带有时间戳的数值型数据，且设备具有静态属性，这些特点决定了时序数据具有较大的同质性和规律性，因此具有较高的压缩比，导致时序数据库的数据压缩能力成为另一个重要的技术方向。

从应用场景来看，时序数据库广泛应用于物联网、工业互联网、数字能源和金融领域。在物联网和工业互联网中，时序数据库用于设备监测、故障诊断和预测性维护，提高设备可靠性和安全性，降低运维成本。在数字能源领域，它助力智能电网、能源管理和能源审计，优化能源消耗，降低成本，支持可再生能源并网和调度。在金融领域，时序数据库用于股票交易分析、风险控制和金融监管，实时监测市场动态，为投资者和金融机构提供决策支持，并保障金融市场稳定和安全。除此之外，时序数据库在车联网、智慧水务、智慧矿山、智慧园区、智能家居和健康监测等领域都有广泛的应用前景。

从国内时序数据库产品的流行度来看，墨天轮中国时序数据库排行榜累计收录 44 款产品，当前排名前三位的分别是 TDengine、KaiwuDB 和 DolphinDB。2023 年，时序专项榜单排行发生较大变化，其中晋升最快的是由浪潮自主研发的分布式多模时序数据库 KaiwuDB。据悉，KaiwuDB 面向 IoT 领域做了定向优化，创新研发“就地计算”核心技术，可轻松实现百万级数据秒级入库，千万级记录查询毫秒级响应，1 秒完

成 20 亿记录数据探索，10 秒完成 500 万记录数据 15 层下钻，切实保障物联网场景海量时序数据的高速入库、极速查询；同步在自适应时序引擎、异构加速计算技术等方面持续发力，已在工业物联网、数字能源、车联网、智慧产业等快速发展的新兴数字产业落地。

2024年1月共44个数据库参与

排行	上月	半年前	名称	模型	数据处理	部署方式	商业模式	专利	论文	案例	资质	书籍	岗位	得分
	1	1	TDengine	时序	-			4	0	0	1	0	0	134.28
	2	↑↑ 4	KaiwuDB	时序				3	0	0	3	0	0	68.71
	3	↓ 2	DolphinDB	时序				3	0	7	1	0	0	64.83
4	↑ 5	↑↑↑ 12	openGemini	时序	-			7	12	0	1	0	0	46.01
5	↓ 4	↓↓ 3	loTDB	时序	-			45	22	0	2	0	0	42.74
6	6	↑↑↑ 9	CnosDB	时序				14	13	0	2	0	0	29.77
7	7	↓↓ 5	CTSDB	时序	-			0	0	0	1	0	0	15.23
8	8	8	YMatrix	时序				1	0	0	1	0	0	10.42
9	9	↑↑↑ 13	NSDB	时序				0	0	0	0	0	0	10.38
10	↑↑↑ 14	↑↑↑ 17	GreptimeDB	时序	-			0	0	0	0	0	0	8.17

图 3：2024 年 1 月墨天轮中国数据库流行度排行（时序数据库）

综上所述，中国时序数据库在工业智能化的发展逐步被业界所重视，处于发展成熟的重要时期，热度持续攀升的同时也有更多的工作待完成。在一段时间内，生态建设将是时序数据库的重点工作。展望未来，笔者认为时序数据库未来将会呈现云边缘计算融合、向 AIIoT 的转变、智能生命周期管理三大发展方向。

（一）云边端计算融合

随着云计算、边缘计算等技术的不断发展，时序数据库将进一步实现云边端计算融合，形成一种高效、智能的数据处理和管理模式。这种模式将利用云计算的存储和计算能力、边缘设备的实时处理能力以及终端设备的感知能力，实现数据在云、边、端之间的无缝流转和处理，更好地满足各种复杂场景下的数据处理需求。

（二）向 AIIoT 的转变

随着 AIIoT 的快速发展,时序数据库将更加紧密地与 AI 技术结合,实现向 AIIoT 的转变。通过与 AI 技术的结合，时序数据库将能够实现对时序数据的自动特征提取、异常检测和预测等功能，提升数据处理的智能化水平，进一步扩展其应用领域和价值。

（三）智能生命周期管理

智能生命周期管理是时序数据库未来发展的重要方向之一。具体来说，智能生命周期管理包括数据的过时淘汰策略、预计算机制、降采样机制和数据压缩存储等方面。通

过这些策略的实施，时序数据库将能够提高数据处理效率、优化存储资源并确保数据的新鲜度和精准度。

11、向量数据库

一、向量数据库是什么

向量数据库是一种专门为管理向量数据而设计的数据库系统。这里的“向量数据”指的是那些通常用于表示非结构化数据的特征，例如图片、音频和视频等。非结构化数据不适合用传统的表格数据库存储和处理，因为它们包含的信息无法有效地以行和列的形式表示。相反，这些数据类型通过机器学习算法提取的特征通常表示为向量，这是一种能表达高维数据空间点的数学结构。

随着人工智能和大模型技术的快速发展，模型在理解数据的语义内容方面变得越来越强大。这促使了向量数据应用场景的扩展，并且使得高效地存储和检索这些向量数据成为了一个重要议题。向量数据库因此诞生，目的是解决这一挑战。

向量数据库的核心功能是理解 and 处理高维数据的相似性。它使用各种机器学习算法，如近邻图、聚类、局部敏感哈希（LSH）等，来执行多种复杂的数据操作。这些操作包括最近邻/最远邻检索（即找到与给定向量最相似/最不相似的向量）、聚类计算（即将相似的数据点分组在一起），以及相似性过滤等功能。

与传统的向量搜索服务和向量检索库相比，向量数据库从一开始就非常重视数据库的关键能力，如数据持久性（Persistence）、一致性（Consistency）、可用性（Availability）、可扩展性（Scalability）和安全性（Security）。之所以称之为向量数据库，是因为开发者希望向量数据的处理可以像处理结构化数据一样高效和易用。

二、向量数据库的发展历史与国内现状

向量数据库的出现与非结构化数据的激增和深度学习技术的快速发展紧密相连。传统数据库面临着处理高维度数据的挑战，特别是在搜索和分析这些数据时。随着深度学习算法越来越多地应用于图像、音频、视频等非结构化数据，产生了大量的向量数据，这些向量数据需要专门的存储和检索机制，从而促成了向量数据库技术的诞生。让我们来简单回顾向量数据库的发展历程：

向量数据库的发展历程：

- 2012 年：深度神经网络的崛起

这一时期标志着深度学习技术的重大突破，其中向量成为了处理非结构化数据（如图像、音频、视频）的重要数据结构。

- **2017 年：Facebook 开源 Faiss 框架**

Faiss 是一个专门为高效相似性搜索和密集向量聚类设计的库。其开源加速了向量检索技术的发展，对于处理大规模数据集特别有用。

- **2018 年：BERT 和其他训练语言模型的出现**

BERT 等模型的出现推动了自然语言处理的巨大进步，也促进了向量数据库技术的发展，这些预训练模型通常是向量数据的主要来源。

- **2019 年：Milvus 正式开源**

Milvus 是全球首款开源的向量数据库，它具有高效进行向量搜索和存储的能力。Milvus 的开源不仅代表了向量数据库技术的正式定义和推广，也象征着这一技术领域的重要发展。

- **2021 年：Milvus 2.0 发布和 Pinecone 服务的推出**

Milvus 2.0 版本的发布将向量数据库带入了云原生时代，提供了更高的性能和可扩展性。同年，Pinecone 推出了其向量检索云服务，进一步推动了这一领域的发展。

- **2023 年：向量数据库作为大模型的记忆体**

随着大型模型的爆火，向量数据库作为大模型的长期记忆体的概念受到广泛关注。RAG, Agent 等 AIGC 应用也成为了向量数据库的重要场景。

总结：

在国内,Zilliz 开源了全球最流行的向量数据库 Milvus,并于 2022 年正式发布 Zilliz cloud,一款公有云全托管的向量检索 SaaS 产品。腾讯,阿里等云厂商也纷纷跟进,发布了自己的向量数据库产品。随着大模型技术的不断落地,向量数据库的需求也会大幅提升。随着越来越多的厂商和开发者不断涌入,向量数据库也在逐渐分化,并且逐渐演进出越来越多与传统数据库不同的能力。

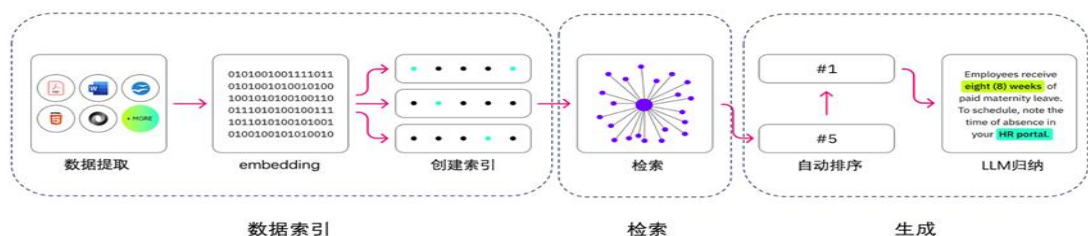
三、向量数据库的典型应用场景

	Content Type	应用场景	搜索		推荐系统	大模型增强	风控	其他
			相似性搜索	多模态搜索				
Content Retriever	Image	Embedding (Vectors)	以图搜图		UGC图片分析		图片侵权 商标查重 涉黄涉恐图片风控	疾病诊断与智能读片
	--Character		OCR					
	--Face		人脸识别					
	--FingerPrint		指纹识别					
	Video		海量视频检索 视频去重	自动驾驶视频 数据检索			视频查重 敏感人物片段定位	
	Audio		声纹识别 ASR语音识别					
	Phrase		知识库检索 搜索词提示 全球专利查询		消费者倾向分析			
	Sentence				社交媒体、用户 评论分析 实时舆情监控	对话机器人	论文查重 敏感内容风控 评论刷分检测	
	Paragraph				新闻推荐 情感分析	专有知识库加强 (提示词工程) 大模型缓存 写作助手		
	User		用户自动分类 (标签)		精准营销 (产 品-用户)		欺诈检测 羊毛党判别 高风险行为判定 用户信用分评估 电商平台商户评级	
	Product/Item		以图搜商品 以文字搜商品		相似/相关产品 推荐			
	Movie/Music				多媒体推荐			
	File						文件查毒 盗版文件监测	
	Chemical Formula							帮助小分子药物定 位蛋白质上的靶点
	Investment Portfolio							投资组合优化

关键场景：

(一) RAG (知识库)

向量数据库支持用户在 LLM 外存储领域特定信息和最新的私有数据。用户提问时，LLM 应用先使用 embedding 模型将用户问题转化为向量。随后，在向量数据库中进行相似性搜索，并获得返回的 top-K 个与用户问题最相似的结果。最后，合并返回结果与原始问题，生成一个包含上下文的新 Prompt，以便 LLM 给出更准确的回答。



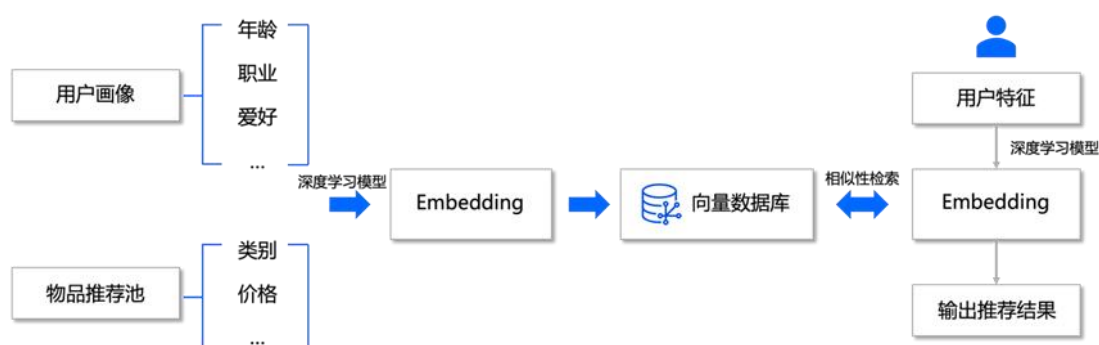
典型用户需求：

- 全链路，覆盖 chunking，embedding，ranking 和向量检索等基本能力
- 对召回质量要求高，同时关注细节信息和上下文召回
- 可以灵活地添加元信息并且低成本地存储原始数据

- 多租户，最大支持千万级别的租户规模
- 存储成本低，单文档成本不超过 1 元每月

（二）推荐系统

使用向量数据库可以搭建快速、灵活、可靠的推荐系统，从而提升用户体验、促进业务增长。一般会先使用深度学习模型将用户购买行为和商品相关数据（例如商品特征、商品描述等）转化为 Embedding 向量。然后使用向量数据库存储和检索 Embedding 向量。向量数据库进行向量相似性搜索，比较和计算用户向量和商品向量之间的距离，从而召回 Top-K 个最相关的结果。最终，推荐系统为用户推荐其可能感兴趣的商品。

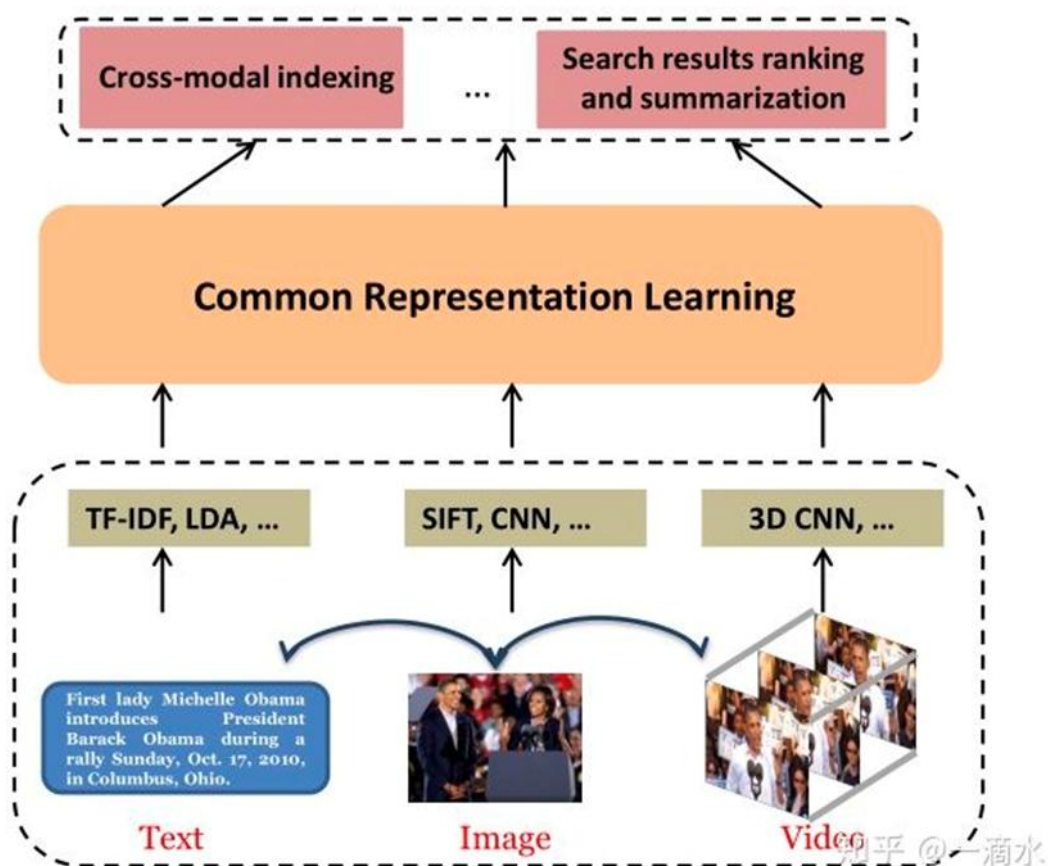


典型用户需求：

- 数据量较小，百万到千万量级
- 高 QPS，延迟要求通常在 20-50ms
- 可用性要求较高，追求 99.99 的 SLA
- 经常有构建全量数据的诉求，支持数据离线导入和 Alias 能力
- 业务流量波动大，需要弹性和自动扩缩容能力

（三）多模态检索

使用向量数据库可以实现多种不同模态的数据（如文本、视频、音频、图片等）进行联合相似性搜索。在一些认证场景中，为了确保准确率，需要对用户的人脸，声纹和指纹多个特征进行联合比对。一般会使用人脸 embedding 模型、声纹 embedding 模型和指纹 embedding 模型将用户的各个特征转化为 embedding 向量并存入向量数据库的多个向量字段中。认证的时候将待认证用户的人脸、声纹和指纹也使用相同的 embedding 模型转成向量再去向量数据库中分别计算各个向量字段的相似度，并做综合打分排序，只有综合得分超过特定阈值才能通过认证。



典型用户需求：

- 搜索质量高，需要多模型或者多向量检索能力
- 延迟敏感，直接服务 Critical 业务场景
- 支持结构化标签的过滤
- 支持丰富的向量检索语意，比如过滤掉部分相似图片，根据图片相似性进行聚合
- 数据量大，扩展性要求高，成本敏感

除上述应用场景外，向量数据库还可以用于文本/语义搜索、图片相似性搜索、音视频相似性搜索、DNA 序列分类、欺诈保护、AI 制药、版权保护、自动驾驶等应用。

四、向量数据库的技术挑战和最新进展

向量数据库在 AIGC（人工智能生成内容）时代确实因其简单的应用和高效的搜索特点而备受关注。然而，这项技术仍面临着许多挑战，这些挑战在一定程度上限制了它的进一步发展和广泛应用：

1. **查询语言：**目前缺乏像 SQL 那样的统一查询语言，使得开发者在使用和优化查询时面临难度。

2. **性能问题**: 在处理大规模非结构化数据时, 检索性能 (QPS/Latency) 可能成为业务发展的主要瓶颈。
3. **高昂的存储成本**: 许多向量数据库依赖于内存索引, 这导致存储成本较高。
4. **查询准确率不足**: 尤其是在处理一些关键词查询时, 单纯的稠密向量检索效果可能无法满足应用需求。
5. **数据安全**: 随着向量数据在各领域的广泛应用, 其数据安全性问题变得日益重要。
6. **容灾能力**: 向量检索业务的实时性要求高, 因此故障恢复的速度非常关键, 目前很多向量的数据都是嵌入式或者单机方案, 并没有关注稳定性和高可用。
7. **多租户支持**: RAG 业务中, 每个用户都有自己的私有知识库, 如何在一个集群中支持百万级别的租户的数据并进行隔离是一个技术挑战。
8. **弹性需求**: AIGC 业务的发展迅速, ChatGPT 获得 2 亿用户只花了 2 个月时间, 这对向量数据库的弹性和扩展能力提出了很高的要求。

五、向量数据库发展趋势

(一) 混合查询能力- 稀疏向量 + 稠密向量混合检索

向量数据库在未来不仅能存储并搜索深度学习模型产生的稠密向量, 还能存储搜索由 BM25/SPLADE 等算法生成的稀疏向量, 通过稀疏向量 + 稠密向量混合检索来整体提升查询准确率。

(二) Serverless

为了解决弹性需求, 使用存算分离架构的 Serverless 方案将是必然趋势。Serverless 可以动态扩缩容满足业务的快速发展, 减少用户的运维成本。

(三) 多级分层存储

为了降低高昂存储成本, 基于内存, 本地磁盘, 对象存储的多级存储结构一定会应用于向量数据库。对于不同热度的数据, 分开放在不同的存储介质, 同时支持基于租户进行冷热隔离。DiskANN 索引和基于 MMap 的换进换出技术将会广泛应用。

(四) 异构计算

随着芯片技术的飞速发展, 向量数据库将会支持 GPU, FPGA, TPU 等定制硬件进行异构加速计算, 一些基于定制硬件的新型向量索引也会诞生。

(五) 智能化

不同向量索引的参数众多，关于参数的设置一直是用户使用的难题。未来的向量数据库需要具备智能化调参的能力，基于 AI 模型智能完成索引和查询参数调整，在查询过程中智能剪枝降低查询代价。

（六）一体化

向量数据库逐步集成 embedding, ranking 等基本能力，帮助用户一站式完成非结构化数据检索，做到 file in file out。

12、流式数据库

一、融合型流式数据库解决了哪些问题

数据已成为现代社会的重要资源之一，对于决策链上的重要数据，数据价值的时效性要求极高。在海量、无穷的数据集中，实时、快速处理新产生的数据，为决策链提供详实、准确、可靠的数据，使其价值最大化，流式处理系统应运而生，比如：Storm、Spark Streaming、Flink、ksqlDB、Materialize 等等。

传统流式处理系统蓬勃发展的同时，其痛点也很明显：

1. **功能单一**，只能做流式计算，没有存储，往往要配合 Hadoop 大数据组件或者数据库来使用。
2. **不易使用**，遇到复杂的处理功能，用户需要编写 Java 或 Scala 代码，难以快速响应复杂多变的业务需求。
3. **难以维护**，现有流式处理系统需要搭配 Hadoop 以及数据库等多个组件满足客户业务需求，技术栈复杂，安装维护困难。

面对上述痛点，目前的技术趋势是将流式处理融入关系型数据库内核，把流式处理和关系型数据库融合在一起，通过一体化（All-In-One）的解决方案和统一的 SQL 使用方式，来帮助用户简化技术栈、实现统一的管理。

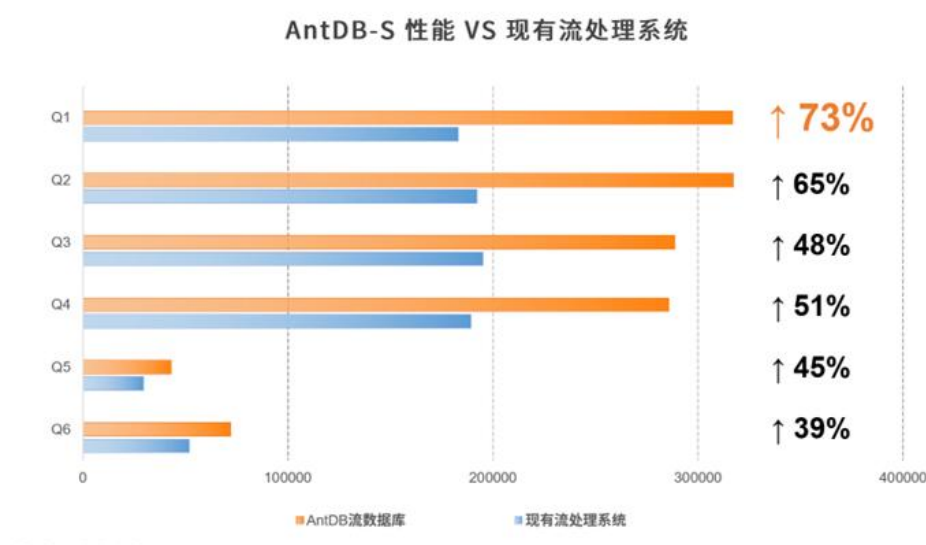
二、融合型流式数据库创新应用案例

流式处理的业务应用范围非常广泛，几乎涉及各行各业。比如：通信、政务、金融、交通、新能源、工业物联网、零售等行业；覆盖的业务类型有：实时监控、预警、数据流转、消息推送、实时特征提取、实时分析、实时报表等；在融合型流式数据库（如：AntDB-S）中这些实时处理业务场景，仅用 SQL 语句就可以完成构建。

以 TPC-C 的零售业务模型为例，在线批发交易业务中包括：客户下订单、支付订单、投递送货等事务操作。这些事务会更新库存表、插入更新订单表、更新客户余额、插入交易历史表，AntDB-S 对库存表、订单表、交易历史表进行流式处理，分别创建为三个流对象，然后通过对库存流对象进行流式过滤查询，来实现对库存的实时预警。对订单流对象进行滚动窗口统计，实现各地市订单情况的实时统计数据视图的展现。对交易历史流对象做滑动窗口统计，用于实时展示销售业绩。

相比传统流式处理系统，融合型流式数据库带来如下好处：

1. **功能全面**：既可以作关系型数据库使用也可以做流式处理系统使用，甚至混合使用比如流表 join、流式 ACID 等；
2. **更易使用**：用 SQL 操作所有流处理功能，SQL 是经过事实证明的最佳数据处理语言，使用方便可以快速响应业务的复杂多变；
3. **维护简单**：与数据库共用一套维护体系，没有多而杂的技术堆栈；
4. **性能提升**：在实验室环境下，同等数据体量可带来平均 50% 的性能提升，具体参见下图。



图：融合型流式数据库性能对比

三、融合型流式数据库展望

融合（All-In-One）是一种趋势，如 HTAP、湖仓一体、流批一体等，未来还将出现更多特性的融合。融合型流式数据库在实现交易型数据库和流式处理引擎的融合后，接下来将进一步融合时序存储、多维向量检索等。

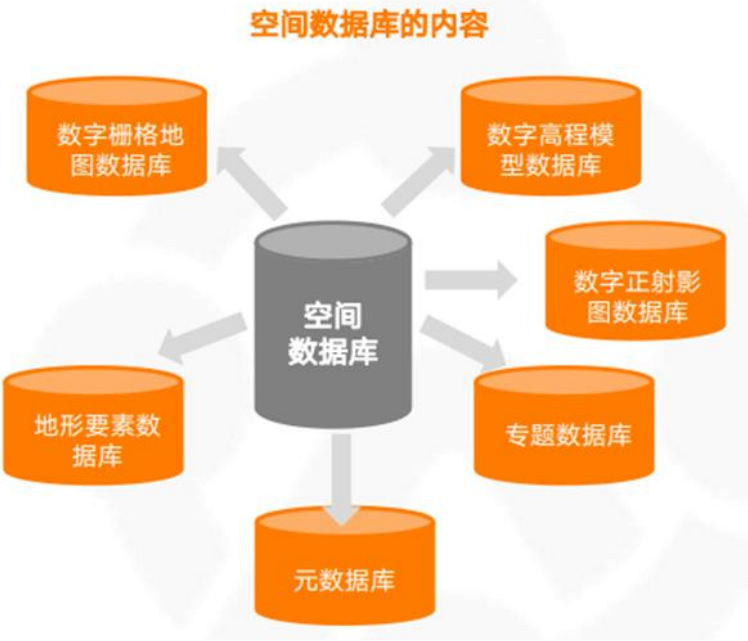
在实际应用中,时序计算和流式计算常常结合使用,以提高数据处理和分析的能力。时序计算可以对事件序列数据进行预处理和特征提取,并将处理后的数据传输给流式处理系统进行实时分析和决策,进一步提高数据处理和分析的能力。

随着人脸识别和 ChatGPT 的普及,向量数据库发展迅速,是专门用于多维特征向量快速检索的数据库。然而,对于复杂业务,还需要其他数据库或大数据组件的支持,学习和维护成本较高。融合型流式数据库计划将 ANN 算法(如 IVF-PQ、hnsw、nsg 等)作为索引,同时实时入库多维特征向量,并使用流式处理引擎构建 ANN 索引,检索时则利用流批一体特性快速匹配特征相似度并进行关联分析。这样可以实现多维特征向量实时入库、快速检索和复杂分析。

13、空间数据库

一、空间数据库简介

空间数据库是一种能够有效地存储、操作和查询空间数据的数据库管理系统。它包括表示在几何空间中定义的对象的空间数据,以及用于查询和分析此类数据的工具。将空间数据与数据库联系起来的三个要素是数据类型、索引和函数,大多数空间数据库都允许表示简单的几何对象,例如点、线和多边形。一些空间数据库处理更复杂的结构,例如 3D 对象、拓扑覆盖、线性网络和不规则三角网络。



二、空间数据库发展状况

国外数据库产品对空间数据管理的支持起步较早，早在 20 世纪 80 年代诞生的关系型数据库 Ingres 就具备了基本的几何对象数据类型管理能力。经过近 40 年的发展，国外关系数据库在空间数据存储和管理方面的研究经历了从二维几何对象存储管理，到支持简单三维几何对象数据存储管理，再到栅格、多媒体等多维数据的管理。

空间数据管理模式主要包括文件系统管理、文件-关系型数据库混合管理、全关系型数据库管理、对象-关系型数据库管理、面向对象数据库管理、关系-非关系数据库并存管理等。目前使用比较广泛的主要有全关系型数据库管理模式和对象关系型管理模式两种。

（一）文件-关系型数据库混合管理方案

用一组文件形式来存储地理空间数据及其拓扑关系，利用通用关系数据库来存储属性数据，通过唯一的标识符来建立它们之间的连接。但该方案难以保证数据的存储和操作的统一。

（二）全关系型数据库管理方案

基于关系模型方式，将图形数据按关系模型组织。图形数据和属性数据统一存储在通用关系数据库中，即将图形文件转成关系存放在目前大部分关系型数据库提供的二进制块中。将图形数据变长部分处理成 Binary Block 字段（多媒体或变长文本）。局限性在于不定长记录造成存储效率的下降；实现 SQL 查询要附加接口，因此它只适用于功能简单的 GIS。

（三）面向对象数据库管理方案

面向对象数据库管理方案面向对象型空间数据库管理系统最适合空间数据的表达和管理。局限性在于不支持 SQL 语言，在通用性上受局限；价格昂贵。

（四）对象关系数据库管理方案

对关系数据库进行扩充，使之能直接存储非结构化的空间数据，形成对象-关系型数据库 GIS 系统。局限性在于空间数据对象不能有用户任意定义，用户使用受一定限制。

当前，国外数据库有部分产品已实现对空间数据的支持，其中具有代表性的空间数据库产品为 Oracle Spatial 和 PostGIS。两者都是基于关系数据内核扩展了对空间数据类型的支持。总体来讲，空间数据模块的代码和团队与内核相对独立，由于空间数据管理模块与内核结合度不紧密，因此对三维空间数据的存储及检索方面支持较弱，目前对三维空间数据的存储管理主要是由空间数据引擎实现。

国产数据库对空间数据管理主要有以下几种：

- 基于 PostgreSQL 改造的国产数据库，通过集成 PostGIS 插件实现空间数据存储管理，PostGIS 插件是基于 GPL 协议开源，存在自主可控风险；
- 基于 MySQL 的开源数据库改造的国产数据库，对空间数据管理的支持相对 Oracle 和 PostGIS 能力较弱；
- 自主研发空间数据管理能力的国产数据库，目前对三维空间数据的支持较弱。

国内空间数据库产品在不同程度支持了二维空间数据管理，其中通过自研方式实现空间数据管理的产品，其功能相对较少，无法全面支撑智慧城市的空间数据管理需求。国内数据库产品对复杂的三维几何对象支持不理想。

三、空间数据库应用及发展前景

现实生活中，大部分的数据库都具有空间属性。目前空间数据库在智慧城市、气象、环境、农业和交通等领域中也有广泛的应用。在地理信息系统中，空间数据库主要用于数据的存储、管理和可视化操作，可以支持多种空间数据分析和处理功能。它可以用于制图、统计分析、空间查询、地图制作和数据共享等工作。此外，在智慧城市中，空间数据库也可用于时空数据分析、交通控制和公共服务等领域，在位置服务和导航、城市规划与管理等方面发挥重要作用。

近年来，国家提出“推动互联网、大数据、人工智能和实体经济深度融合，建设数字中国、智慧社会”，智慧城市建设已成为推进城市治理体系和治理能力现代化重要抓手。随着智慧城市、数字孪生、5G、物联网等领域的深入推进，交通出行、定位导航、社交活动、推荐系统等各种应用层出不穷，产生了海量的非结构化数据，空间大数据量级呈指数型增长，且数据类型多样、来源广泛、信息量大。

四、挑战与解决方案

面对多源异构的海量空间数据，承载其运行的底层空间数据库的性能直接关系到数字孪生 BIM/CIM 平台效力的发挥。在国产信创环境下，空间大数据存储与计算面临较大的挑战：

一是非结构化的空间数据缺乏统一的标准与规范。海量的物联数据、轨迹数据、视频数据、BIM 数据、政务数据等城市地上下和室内外的时空数据，在数学基础、空间基准、尺度、现势性、编码和格式等方面缺乏统一的标准和规范，阻碍数据接入和融合处理，严重影响数字孪生城市发展。

二是传统云计算处理模式难以保证实时性、可靠性和高效率。针对城市级、大规模和高速变化的城市政务管理和运行感知等动静态数据，特别是巨量的 BIM、倾斜摄影和时空轨迹等非结构化数据，传统云计算的处理模式在实时性、高效性和可靠性等方面存在严重不足，亟须研究集存储、索引、搜索和计算为一体化的分布式存储和计算架构，研发高性能 CIM 搜索和计算引擎，从底层支撑城市时空大数据的分析、计算、查询等性能。

以深圳 CIM 平台测试评估和调研分析情况为例，国产信创环境下，支持海量非关系型数据的国产数据库系统性能还不足以支撑 BIM/CIM 平台的建设与应用。同时，空间数据管理对数据库也有一大挑战，如结构复杂、检索方式不同、不同数据库管理方式不同难以兼容等。智慧城市建设过程中涉及大量的公共数据和关键信息，打造国产化自主可控的高性能空间数据库系统势在必行。

在此背景下，YashanDB 推出了 YashanDB for GIS。在最新的 YashanDB V23.1 版本里，已支持空间数据管理能力，原生空间数据类型利用 UDT(自定义类型)能力进行扩展，用于存储和访问符合开放地理空间信息联盟（简称 OGC）制定的 SFA SQL 标准的几何对象。

YashanDB 在城市级范围 55W 栋建筑物数据集上进行了对比测试，在常用空间计算函数上性能平均优于开源 3 倍，为 CIM 平台提供了更强的空间对象计算能力支撑。



图：YashanDB for GIS 能力

五、技术趋势

从 YashanDB 目前的实践总结，空间数据库主要有以下几个方向值得关注。

1.GIS 与其他类型数据融合，在城市数字孪生场景中需要将细粒度建筑模型 BIM 与 GIS 底板融合，实现高效跨模型数据计算；空间数据也需要考虑其动态变化情况，需要与 IOT 数据联合处理。

2. **空间计算实现分布式并行化**，面对空间数据规模的持续增长，传统基于集中式的计算能力容易形成瓶颈，通过分布式实现空间计算的可扩展。
3. **空间数据安全**，如何在数据的应用和共享中避免关键敏感数据泄露。

14、数据仓库

实时数据仓库，顾名思义，其最大的特点就是“实时”，它以其高效的数据处理能力和即时的数据分析能力，成为现代企业不可或缺的数据基础设施。

一、数据仓库的发展历程

纵观数据仓库的发展历程，数据仓库的演进经历了三个阶段。第一阶段，即在 2010 年之前，传统数据仓库解决方案包括 Teradata、Greenplum、IBM Netezza 来处理数据分析问题。随着大数据时代的来临，数据的产生速度呈现爆炸式增长，传统的数据仓库已经无法满足企业对大规模数据处理和分析的需求，因此第二阶段约在 2010 年左右，随着谷歌三驾马车的问世，人们逐步开始使用 Hadoop 大数据平台进行海量数据的离线处理和分析。在传统的数据仓库中，数据的采集、处理和分析往往需要经过一系列的批处理过程，这导致了数据的延迟，使得决策者无法及时获取最新的数据情报。

如今已经进入了第三阶段，现代化的数据仓库产品开始涌现，这些产品结合了传统数据仓库的可靠性和性能优势，更重要的是，**具备了对大数据的高效处理和实时分析能力**，通过一系列先进技术，实现了数据的实时采集、实时处理和实时分析，大幅缩短了数据从产生到被利用的时间周期，帮助实时数据的业务价值得到充分发挥。

二、传统数仓的痛点

面对云时代和大模型时代的到来，传统数仓存在的瓶颈带来了多种用户的痛点，其中包括：

（一）缺乏弹性

传统数据仓库的计算和存储是紧密耦合的，计算资源和存储资源按某一比例强绑定，因此用户在扩容时，必须同时扩容计算资源和存储资源，在扩缩容、运维、迁移上都存在一定的挑战。当企业遇到负载高峰时刻或需要紧急得到某个报表结果时，传统数据仓库无法及时扩资源，导致大数据系统无法弹性、快速地分析业务数据，错失了充分挖掘数据价值所带来的商业机会。

（二）成本高昂

传统数据库价格高昂的软硬件、开发运维人员的高昂薪资需要企业进行巨大的前期投入。传统数据仓库客户的生产环境资源利用率，无论是存储或是计算资源往往都不尽如人意。随着存储和工作负载需求的日益增长，面临数据库的扩容和升级时，由于传统数据库架构存储和计算的紧密耦合，往往需要企业花费巨大的运维和时间成本，且操作繁琐。

（三）木桶效应

传统 MPP 数据库架构存在“木桶效应”，数据库实例集群整体执行速度取决于最“短板的”节点的性能。因此，一个节点的表现往往会“拖垮”整个数据库实例集群的性能，导致查询速度变慢。随着时间的推移，业务的增长，企业往往需要在 1-2 年后对数据库实例集群增加计算节点，此时，无论新的计算节点性能如何好，数据库实例集群总体性能都会受制于老的节点。因此真实生产环境中，我们常常见到客户在需要扩容时，采取重新新建数据库实例集群的方式。

（四）数据孤岛

随着业务的发展，数据量的增加，和信息化建设的需求，企业会为不同部门建设相应的业务信息化系统。我们在真实客户场景中，常常看到很多企业有成百上千个数据库实例集群，但这些数据库实例集群的元数据往往都是一样的。这种情况下，很多元数据会在不同数据库实例集群间存在不一致的版本信息。此外，如果企业需要做跨数据库实例集群的访问，往往非常困难，会造成数据孤岛的存在。

（五）运维成本

对于传统 MPP 数据库，企业往往会需要配备运维人力，且对运维、开发人员要求高，需要相关人员掌握复杂的技术栈，技术的更新迭代迅速，相关人员需保持积极的知识更新意识。相关人才市场较小，人才匮乏。高昂的学习成本造成用户使用过程中性能差、故障率高、故障修复时间长等问题；

三、数据仓库三大现代化趋势

在数据仓库技术发展的过程中，旺盛的市场需求催生着新的概念、技术和产品不断涌现，总体而言，数据仓库呈现三大现代化趋势——**实时化、统一化以及云原生化**。

（一）实时化——从批量到实时

过去，大多数企业使用的传统数据仓库/大数据平台主要是对历史数据进行批量分析，如果能对数据进行实时分析并将分析结果实时运用到业务之中，毫无疑问将会进一步利用好数据的实时价值并驱动业务进步。因此到如今时代，数据分析逐渐从原来的批量处理演变到现在的实时处理。

以业务分析需求的变化为例，越来越多的企业开始采用实时报表和实时仪表盘展示数据，取代了传统跑批任务生成的报表。而从批量生成的静态报表到交互式分析也是另一个典型趋势，过去我们只需要跑一份静态报表，而现如今现在很多公司内部都有大量的数据分析师，需要与系统进行快速互动实时产出分析结果。此外，数据结果不再仅限于人使用，逐渐转向为机器和算法使用的实时决策系统。这些变化清晰地展现了一个新的趋势：数据从批量处理逐渐转向实时分析已成为必然。

与此同时，过去数据分析系统主要是给内部的运营决策或数据统计来使用，而随着业务的发展、数字化转型的深化，越来越多的数据分析开始面向业务的外部客户，主要场景包括广告营销报表、物流实时看板、保险客户分析和交易明细查询等。这些都是数据分析需求由内到外的转换，这种转变也要求数据仓库能够适应更多样化的业务场景。

在应对大规模数据的实时分析时，核心挑战来自两个方面：

- 随着数据实时写入数据库，面临的挑战之一是如何以更低的延迟提供数据，需要降低数据传输和处理的延迟，以提高数据的新鲜度，并及时处理最新数据的变化。
- 对于上层数据应用而言，如何提供更快的查询、降低查询耗时，需要持续优化查询性能，提高查询的快速响应度，以满足上层数据应用的性能需求。

因此，在数据实时写入方面，数据仓库的存储引擎需要提供秒级的数据实时更新、在主键表和非主键表上支持高效的实时更新和追加。同时支持数据库的 CDC（变更数据捕获）功能以及 Kafka 的流式数据同步功能，能够实现秒级的数据同步。不止数据可以实时写入和更新，对于表的模式（Schema）也需要进行快速变更，以适应当今快速变化的业务环境，且 Schema 修改期间完全不影响在线业务的运行。

在数据查询方面，数据仓库需要支持多种不同查询负载的高性能查询，无论是高并发点查询、大宽表查询、多表关联查询还是增量库内 ELT，需要具备在一套架构下同时满足高吞吐的 OLAP 分析和 ETL 处理以及高并发的 Data Serving 在线服务的能力，简化在混合工作负载下的技术架构，为用户提供了多场景下的统一分析体验。

（二）统一化——OLAP 技术栈的统一、湖和仓的融合统一

在大数据领域，存在众多的系统和组件，它们往往在架构中扮演着不同的角色。而随着时代的进步，架构“减负”已成为企业发展的重要目标，因此像融合数据库、超融合数据库、湖仓一体、流批一体等具有“融合统一”特征的数据库开始涌现。

结合实际发展来看，过去几年，事务数据库已呈现明显的融合趋势，像 OceanBase、TiDB 以及 PolarDB 等新一代 NewSQL 事务数据库已经出现。同样地，数据仓库走到今天，统一化也是必经之路。OLAP 领域拥有大量针对特定场景的工具和技术，企业通常需要部署大量的产品来满足不同分析场景的需求，各种 OLAP 系

统层出不穷，而这带来了组件过多、运维成本高、数据链路长、数据重复存储等一系列问题。

因此 OLAP 系统在朝着统一化的方向去发展，单个系统支持多种数据源、支持多种数据类型、支持多种数据分析负载，在一套系统中减少繁重架构带来的冗余存储和运维压力。All-In-One 的数据仓库平台，更加易于使用和管理，让企业精力从管理复杂的数据基础设施转为关注上层的数据应用。

另一方面，数据仓库在性能方面表现出色，而数据湖则以其开放性和能够存储各种数据的优势而受到青睐。然而无论湖或仓在场景上都具备一定的局限性，因此如今我们正处于数据湖和数据仓库融合的阶段，要想充分利用数据仓库的高性能和数据湖的开放性，整合这两者变得至关重要。

因此对于现代化的数据仓库，湖仓统一融合的最重要特性有两个方面：

- **Federated Query Engine:** 作为一个联邦查询引擎，数据仓库可以访问各种数据湖上的表格式，包括存储在 HDFS、S3 上的 CSV、Parquet、JSON 等文件，以及查询 Iceberg、Hudi、Delta Lake 等数据湖中的数据。
- **Open Data Lake:** 数据仓库还可以作为一个开放的数据湖，供外部查询引擎或者数据科学工具进行读取，这种能力在数据湖与数据仓库融合的背景下，成为数据仓库需要展现的一个重要特点。

以 Apache Doris 为例，作为一个开源实时数据仓库同样也是高效的联邦查询引擎，Apache Doris 可以通过创建数据目录的方式与外部数据源进行映射，例如可将 Hive、Elasticsearch、Iceberg 等数据源映射为外部表，将自动更新源数据、并自动进行外部数据的高速缓存。同时数据目录的更新操作是按需的，可以指定要查询的库和表进行更新，也可以利用插入语句将查询结果插入到内部表中，这些操作仅需一条命令即可完成，具有更高的易用性。

在当前的数据湖与数据仓库的融合趋势中，大多数通常只关注了前一方面，而数据仓库的格式开放性很少被提及。在当今这个数据类型多样且数据负载庞大的时代，数据仓库可能也将面临数据科学或其他形式的大规模分布式计算，如果不开放数据格式，那么 Spark、Python 之类的工具将无法使用。因此在开放数据格式方面，Apache Doris 采用了 HTTP Data API，可以使客户端以并发方式与多个 BE 进行读取，并提供更高的数据读取能力。无论是使用 Flink Connector、Spark Connector，还是通过 Python SDK（数据科学、机器学习）都可以快速访问。

因此数据仓库也必须具备高效可拓展的读取方式，便于数据被外部数据科学或机器学习应用读取、与整个 AI 和数据科学生态进行良好的整合，这也是未来的重要发展方向。

（三）云原生——基于云做好存算分离架构下的弹性计算

新兴云计算基础设施的成熟催生了新的需求，无论是公有云、私有云以及基于 K8s 的容器平台，数据仓库的云原生能带来的核心价值有四个方面：

- 首先是存算分离。在云上提供了高质量、低成本的对象存储系统以及 HDFS 等共享存储系统，将大量的历史数据或冷数据迁移到低成本的存储介质会为企业带来巨大的成本节约。
- 其次，存算分离的特性使得计算变得更加弹性，在业务波峰波谷效应明显的场景中，也可以通过计算资源的弹性调度更好地应对变化。
- 同时，由于进行了存算分离，计算可以实现更好的负载隔离，完全可以根据业务需求进行隔离或者进行读写分离。
- 另外，由于可以共享同一份存储，数据共享变得更加简单。这意味着数据不再需要繁重的迁移工作，同一份数据可以被多个计算业务共享使用。

在数据仓库的云原生趋势下，核心在于如何更好地基于云做好存算分离架构下的弹性计算，在此许多云原生数据仓库进行了许多设计和思考，在此以云原生数据仓库 SelectDB Cloud 的实现为例：

1. 基于共享存储的主数据存储

在存算一体的架构下，数据主要存储在计算节点上，即使使用了冷热数据分层，热数据依旧只在计算节点上存储，计算节点需要依靠自身的多副本机制保证数据的可靠性。在存算分离架构下，计算节点不再存储主数据，而是将共享存储层作为统一的数据主存储空间，这将带来如下收益：

- 上层的计算节点可以做到无状态，可以实现完全关机
- 更便捷的数据共享，不同的集群之间以及不同的仓库可以便捷地进行数据共享
- 更简易的数据备份与恢复，以及实现数据的 Time Travel

当然，成熟稳定的 HDFS/对象存储还为系统带来极低的存储成本和极高的数据可靠性，并且大大简化上层计算节点的实现复杂度。

2. 基于本地高速缓存的性能优化

存算分离依赖从网络上读取存储系统的数据来进行计算，在一定程度上会造成计算性能的下降，这也是相较于存算一体架构的主要劣势。为了解决这一问题，可以在本地利用 SSD 提供高速缓存。

正如存算一体通过冷热数据分层技术来大大缓解了存储和计算必须同时扩展的问题，同样在存算分离架构中引入计算节点本地高速缓存实际也是融合了存算一体的能力。这种本地高速缓存加上共享存储系统，我们也可以称之为混合模式，无论是 Snowflake 还是 RedShift，实际上都是采用了这种方式来应对底层对象存储系统性能不佳和网络传输带来的性能下降。

引入本地高速缓存后，系统会自动根据 LRU 来缓存最新写入和访问数据，当然也可以手动设定表的缓存策略。由于只是缓存，因此本地只存储了单个副本，这样大幅提升了缓存利用率，相比存算一体模式可以降低 2/3 的高速存储使用。

另外，在存算一体的模式下，多数据副本冗余存储、在三副本上都需要独立进行数据合并计算。而在存算分离下，只有一个节点进行数据合并计算，这样就可以降低 2/3 的数据合并计算量。所以，通过引入本地高速缓存，不仅仅可以基本达到原来存算一体的性能，在有些情况下还会超越原来存算一体的性能。

3.多计算集群实现工作负载隔离

用户通常希望对同一份数据上的分析负载进行隔离。例如，导入的工作负载与查询的负载进行隔离，Adhoc 的大查询负载和在线点查询的负载间相互隔离，避免不同负载间相互资源抢占。

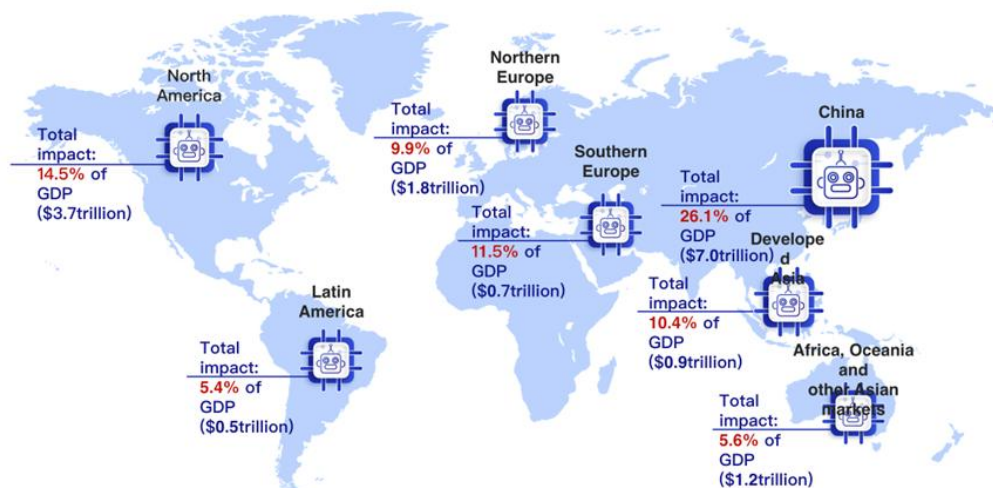
在存算分离模式下，提供了同一个仓库多个物理计算集群的隔离方式。因为主数据存储在共享的对象存储上，因此用户可以按需创建多个计算集群但共享同一份数据。计算集群之间是物理隔离的，可以独立扩缩容，其计算节点的本地高速缓存都是隔离的，这样保证了尽可能比较好的隔离性。

4.计算节点弹性扩缩容

云原生架构强调弹性，数据存储 HDFS 或对象存储上，可以按需扩缩容。每个计算集群的计算节点，能够根据负载需求快速地进行扩容或缩容，也可以根据特定时间进行自动扩缩容。此外，还需要支持集群的自动启停，当没有负载时会自动停止，有查询负载时会自动启动，通过弹性计算的能力，可以进一步节省计算成本。

四、AI 时代下新一代数据仓库

随着人工智能（Artificial Intelligence）技术的发展，海量的数据被广泛收集、存储和利用。在这个数字化时代，AI 大模型成为处理和分析这些数据的核心，而人类正经历着前所未有的模型数量和规模上的指数级增长。



Source: PwC 全球人工智能报告

AI 大模型带来了前所未有的机遇，众多机构均认为：AI 将引领下一波全球 GDP 的增长。根据麦肯锡 2023 年 6 月报告¹，生成式 AI（基于大模型）每年会为全球 GDP 贡献约 2.6 万亿至 4.4 万亿美元，相当于英国 2021 年 GDP 总值（3.1 万亿美元）。高盛也在 2023 年 4 月报告²中指出，生成式 AI 可以为全球 GDP 贡献 7% 的增长。此外，早在 2017 年，普华永道的报告³就曾指出，至 2030 年，AI 对中国 GDP 的贡献可达 7 万亿美元，相当于中国 GDP 的 26.1%；

为了助力企业优化计算瓶颈，充分利用和发挥数据规模优势，构建核心技术壁垒，更好地赋能业务发展，新一代数据仓库需要能够整合企业所有多模态数据资源，提供多模态大模型下数据计算支撑，更贴近数据科学家的需求和使用。

（一）数据仓库上云虚拟化的核心价值

新一代数据仓库需要能够采用领先的数仓虚拟化技术，将多个数仓统一整合到一个高可用的云虚拟数仓，打通多云的数据管道，数据计算资源按需扩缩容，提升数仓的敏捷性和弹性，助力企业降低数仓管理复杂度，具有可扩展性、灵活性和可靠性等优点。典型产品包括拓数派的云原生虚拟数仓 PieCloudDB 产品。

1.降低数仓硬件和管理成本

通过云原生存算分离架构，物理数仓被整合到云原生数据计算平台，支持根据数据授权动态创建虚拟数仓，打破数据孤岛，解决数据多副本问题，帮助企业降低数仓管理复杂度，以更低的成本实现存算资源在云上更灵活的配置。

2.提升数据计算资源利用效益

数据计算资源按需扩缩容，实现计算资源配置最优化，提升数仓的敏捷性和弹性，打开无限数据计算空间，支撑更大模型所需的数据和计算。更好地赋能业务发展并走向绿色。

新一代数据仓库天然支持云环境，无需进行额外的定制。企业可根据对资源的需求，灵活地以低成本和高效的方式，单独地进行存储或计算资源的弹性扩展，提高了资源的利用率，节省空间成本和能耗开销。并支持随着负载的变化实现高效的伸缩，轻松应对 PB 级海量数据，具有高容错性、易于管理和便于观察等特性。结合可靠的自动化手段，能够轻松地对系统做出频繁和可预测的重大变更。

云原生的“即开即用”特性为企业节省了大量运维开支。由于其计算节点部署于云端，摆脱了物理限制和潜在的延迟，可随时随地通过互联网轻松管理，无需任何硬件。数据随时随地可用，无需处理任何后端技术问题，为企业进行跨部门、跨区域的数据共享和协作开辟了捷径，保证了企业的全球化进程。

3.高在线、高安全、高可靠

数据仓库上云，用户最关注的一点就是数据安全问题。传统数据平台将文件和资源存储在同一主机中，以主备节点数据方式补偿节点宕机时间，严重影响数据时效性，增加了运维的成本和难度。云时代下的数据仓库通过数据透明加密（TDE）等技术保证所有数据在落盘前完成加密，服务器无感知技术（Serverless）利用云上无限计算资源和弹性保证了计算永远在线可用，S3 存储和跨云灾备能力保证了永不丢数，确保业务连续性。

（二）数据仓库上云虚拟化的技术突破

为了实现数据仓库上云虚拟化，为企业提供全新基于云数仓数字化解决方案，助力企业建立以数据资产为核心的竞争壁垒，以云资源最优化配置实现无限数据计算可能，新一代数据仓库需要实现以下技术突破。

1.打造云原生存算分离架构，云上资源弹性分配

云时代下的数据仓库在设计与实现的过程中，需要充分考虑云平台的弹性和分布式特性，实现元数据、用户数据和计算资源的解耦分离。通过把元数据，用户数据和计算资源进行解耦，元数据可以被当作保险箱数字 Key 数据，用户数据可以被当作保险箱里的数据黄金。用户只需要对数字 Key 进行交换，就可以访问保险箱里面的数据。把云上无感知计算当作一堆计算器，需要的时候，根据授权的数字 Key，拉到对应的保险箱数据即可进行计算。

新一代的数据仓库需采用高效并行的方式进行数据加载和处理，处理速度随节点增加而提升，支持流数据快速加载。通过云原生虚拟数仓的存算分离架构，实现多集群并

发执行任务,企业可灵活进行扩缩容,随着负载的变化实现高效的伸缩,轻松应对 PB 级海量数据,具有高容错性、易于管理和便于观察等特性。并结合可靠的自动化手段,轻松地对系统做出频繁和可预测的重大变更。

2.打造云原生交互存储与缓存架构设计

对象存储天然适应云原生环境,与云计算平台、容器编排技术等其他云原生技术无缝集成。此外,相对于传统的存储方式,对象存储通常具有更低的成本,成为数据仓库更有优势的一种存储选择。然而,由于对象存储通常是基于分布式系统和网络存储的,数据的传输和检索通常需要网络进行,因此相对于本地存储,会存在一定的延迟。这一缺陷也对云时代的数据仓库底层的存储和存储引擎也提出了新的挑战。

新一代数据仓库需要能够具备强大的存储适配接口能力,确保支持各种类型存储,保证和不同云环境的兼容性。此外,各个计算节点需针对元数据和用户数据均设计多层缓存结构,避免网络延迟和数据移动,提高计算效率,保证用户的实时性需求。针对底层对象存储,新一代数据仓库需设计高效的文件格式,在节省网络请求的同时提高计算效率。

3.云原生优化器的设计与打造

优化器作为数据库管理系统中的关键技术,对数据库性能和效率具有重要影响。针对云原生和分布式场景,优化器需要实现包括聚集下推、预计算、Block Skipping 等高级特性,全面满足各种复杂的分析查询需求。

4.向量化执行器的设计与打造

数据分析和应用的重要性日益增长,对于数据平台来说,极致的性能是关键需求之一。为实现更高效的数据并行计算,新一代数据仓库优秀的执行器需要能够充分利用硬件资源,如 CPU 的并行计算能力和 SIMD 指令集,充分利用了数据并行计算的优势,通过将多个数据元素打包成向量,并同时对其执行相同的操作,提高了计算效率和吞吐量。

5.大模型时代 AI 的支持与集成

新一代数据仓库是大模型时代的分析型数据库升维。对于大模型而言,模型所需的数据都经过了向量化过程,经过向量化的数据可以大幅提升模型的查找效率,降低训练成本。大模型时代下的数据仓库需要能够进一步实现海量向量数据存储与高效查询,助力多模态大模型 AI 应用,支持和配合大模型的 Embeddings,提供对向量的高效存储、索引和查询功能,具备高效存储和检索向量数据、相似性搜索、向量索引、向量聚类 and

分类、高性能并行计算、强大可扩展性和容错性等特性，帮助基础模型在场景 AI 的快速适配和二次开发。

6.对新硬件的支持和兼容

为了加速大数据处理和计算的性能，云时代下新一代数据仓库需要能够充分依赖新的硬件来进行异步计算，例如 GPU、FPGA 等。通过充分利用新一代硬件加速器，数据仓库可以实现更高的计算性能、更低的延迟和更好的扩展性。这将使得大数据处理和计算变得更加高效和可靠，推动云时代下数据分析和决策支持能力的进一步发展。

五、未来展望

未来，随着技术的不断进步和业务需求的持续演进，数据仓库还将朝着更加实时化、统一化和云原生化的方向发展。而人工智能和机器学习技术的应用未来也将会提升实时数据仓库的分析智能化水平，为企业带来更加精准和高效的决策支持。

15、数据湖

一、数据湖的历史

（一）数据湖 1.0

Data Lake 的概念最初由 Pentaho 的创始人兼前首席技术官 James Dixon 于 2010 年提出。

Data Lake 被设想为一个集中存储区，用于保存来自多个源的原始、未结构化、半结构化和结构化数据，而没有预定义的模式。这种设计旨在保存可能具有价值的数
据，并允许数据科学家挖掘和分析大量的大数据。他提出数据湖的概念主要是针对当时数据集市（Data mart）存在的信息隔离问题。

初期的 Data Lake 具有批处理导向的处理方式，这在一定程度上限制了其功能。例如，Hadoop 1.0 版本的 Data Lake 就是以批处理为主，用户需要具备 Java 和 MapReduce 等高级工具的专业知识来与 Data Lake 进行交互。随着 Hive 的推出，人们开始喜欢使用通用的 SQL 来对大数据进行分析，而不是编写复杂低效的 MapReduce 任务。从这一点可以看出，数据湖和数据库的融合是历史的必然。

Hadoop 的推出核心理念是降低在大规模数据下的可扩展性问题，在廉价的硬件上通过分布式计算来降低专有硬件的门槛。然而，Hadoop 生态继承了从 GFS 传承的设计理念，即用追加方式简化存储引擎的复杂性，因此第一代数据湖缺失了传统数据库最关键的更新能力。在这种范式下，数据库中复杂的存储过程也被简化，或者说劣化成了一种新的批处理方式。

随着越来越多互联网公司采用 Hadoop 作为他们的大数据处理方案，大家开始用 insert (overwrite) 和临时表方式来实现复杂的 ETL 加工，替代存储过程。在 Spark 诞生后，通过 RDD (Resilient Distributed Dataset, 弹性分布式数据集) 的抽象，实现了对大规模数据更高效的内存计算和逐阶段的计算容错机制，从而使数据湖的批处理能力达到了更高的水平。

数据湖 1.0 时代也出现了一批商业化公司，其中 Cloudera、MapR 和 Hortonworks 被认为是"Hadoop 三驾马车"。这三家公司在推动 Hadoop 技术的发展和应用方面发挥了重要作用。尽管开源技术发展迅猛，但 2010 年至 2020 年也是云计算和云原生数据仓库快速发展的十年。

一方面，AWS 等云厂商纷纷推出自己的 EMR 云服务和 Redshift/Bigquery 等云数据仓库服务，而固守在私有化场景下的 Cloudera 并没有完成良好的商业化发展。与此同时，以 Snowflake 为首的云数据仓库产品以其易用性脱颖而出，它们对 SQL 的更全面支持使得大数据开发变得更加平民化。

Cloudera 和 Hortonworks 在 2019 年合并，MapR 也因营收不佳在 2019 年被惠普收购。Cloudera 自 2017 年上市以来一直面临盈利能力和竞争对手压力的挑战，最终于 2021 年 10 月私有化退市。

这三家公司的兴起、合并和转型反映了整个大数据行业 Hadoop 的发展趋势。

(二) 数据湖 2.0

2017 年，Linkedin 开源了 Hudi，2018 年，Netflix 开源了 Iceberg，并将他们捐献给了 apache 基金会。

这两个项目在 2020 年成为顶级项目，它们的核心思想不谋而合。

例如，Hudi 在原有数据湖的基础上支持原地更新 (upserts)、变更捕获 (change capture) 和时间旅行 (time travel) 等功能，使得在大规模数据集上进行高效的数据增删改查 (CRUD) 操作成为可能。

Iceberg 设计了一种新的表格模型，支持复杂的嵌套数据结构，提供了全面的元数据管理，并且能够处理大规模的数据集和频繁的更新。同时，它还设计了更好的元数据

管理机制，通过一个独立的 SDK 暴露接口，使上层计算引擎更容易接入这些全新的表格式。

数据湖在 2.0 时代的查询引擎也逐步发展并丰富。除了最早的 Spark, Flink 在流计算上持续演进，保持行业领先地位。而 Trino/Presto/StarRocks/Doris 等 MPP 计算引擎也在完善交互式分析能力的基础上开发出一系列新功能，使得流计算、批计算和交互式分析有了丰富的选择。

二、湖仓（Lakehouse）的由来

在 2020 年的一份名为《Lakehouse: A New Generation of Open Platforms that Unify Data Warehousing and Advanced Analytics》的白皮书中，Databricks 的创始人提出了一个数据湖仓架构的愿景，该架构结合了数据仓库和数据湖的功能，以实现数据仓库的性能和数据湖的经济规模。这种统一的方法消除了传统上分离和隔离数据的数据孤岛，使 Lakehouse 平台具有出色的性能，并且易于集成和使用。

Lakehouse 通过开放格式的低成本云存储实现事务支持、模式执行和治理、BI 支持、存储与计算的分离等特性，能够适应不同的数据应用，包括 SQL 分析、数据科学和机器学习，并提供数据版本管理、治理、安全性和 ACID 属性。

Data Lake 和 Data Warehouse 结合：Data Lakehouse 的诞生，更开放，可以跟 AI 结合更好。

Data Lakehouse 的出现是为了解决 Data Lake 和 Data Warehouse 的局限性。它通过在 Data Lake 的低成本存储上实施类似于 Data Warehouse 的数据结构和管理特性，实现了数据架构的新类别。Data Lakehouse 结构通过使用开放格式（如 Apache Parquet）、为数据科学和机器学习提供原生类支持，以及在低成本存储上提供最佳性能和可靠性，解决了当前数据架构的关键挑战。

随着企业中数据种类的日益增加，例如文本、物联网数据、图像、音频、视频等，Data Warehouse 的局限性开始显现。此外，机器学习（ML）和人工智能（AI）的兴起也带来了对直接数据访问和非 SQL 基础的迭代算法的需求。Data Lakehouse 结构通过使用开放格式（如 Apache Parquet）、为数据科学和机器学习提供原生类支持，以及在低成本存储上提供最佳性能和可靠性，解决了当前数据架构的关键挑战。

三、数据湖存储的 2023 发展

2023 年的数据湖存储有一些创新和进步的趋势，先来看看几个数据湖格式的发展动态。

（一）Apache Iceberg Updates in 2023

Iceberg 在 2023 年发布了 1.2 1.3 1.4 等几个大版本, 其中最重要的是 1.4 版本:

Apache Iceberg 1.4.0 将默认格式升级到 v2, 此举旨在解决 v1 中存在的元数据问题, 并支持行级删除文件, 以减少写放大现象。

这种格式在两年前就被采用, 并得到了广泛支持, 使得大多数部署能够更平滑地过渡。此外, v2 的能力使得数据操作更加高效, 例如, 对于 Spark 的 MERGE 和 UPDATE 命令采用合并读取策略, 从而提高了整体性能和数据管理实践。

此外, Iceberg 1.4 版本还进一步完善了 Rest catalog 协议, 增加了多表提交 (multi-table commit)、延迟快照加载 (lazy snapshot loading) 等功能, 为后续的多表事务等工作做好了准备。

同时, Iceberg 在 roadmap 中放出了 V3 的设计, 其中主要包括, 对加密的支持 (近期已经实现了部分对 parquet 和 avor 的加密), 对默认值的支持, 以及对相对路径的支持。

另外, Iceberg 在 2023 年大大强化了 pylceberg 的能力, 连续发布了多个版本, 现在 pylceberg 基本具备了 Java 版本的完整功能。Iceberg 在北美的顶级互联网公司中得到了广泛应用, 同时 Snowflake、BigQuery 和 Dremio 等也将 Iceberg 作为首选的湖格式进行支持, 因此它拥有最广泛的生态支持。与其他湖格式不同的是, Iceberg 在一些案例中完整替代了 Hive, 成为整体湖格式方案的解决方案。因此, 在当前人工智能飞速发展的时代背景下, Iceberg 还重点发展了 Python API。

（二）Apache Hudi Updates in 2023

Hudi 在 2023 年发布了 0.13、0.14 两个大版本, 同时也发布了 1.0.0-beta1 的预览。

在 0.13 版本中, 引入了元数据管理服务 Meta Server, 用于简化大量数据表的管理。此外, 还正式引入了 CDC (Change Data Capture), 允许用户跟踪单次提交中的记录变更, 并支持流式源使用。同时, 0.13 版本还增加了对部分列更新的支持。

在 0.14 版本中, 最主要的功能是引入了记录级索引 (Record level index), 提高了大型表的写入性能。记录级索引可以有效替代其他全局索引 (如 Global_bloom、Global_Simple 或 Hbase), 并且可以将索引查找性能提高 4 到 10 倍。此外, 0.14 版本还支持表自动生成键, 对于日志类型的表不再需要为 Hudi 表配置主键, 提高了易用性和灵活性。

1.0 版本是 Apache Hudi 的第一个重要版本，其中最大的变化是表格式（table format）的升级，在 Hudi 1.0 的 spec 文档中，我们看到升级后的格式增强了未来的兼容性、提高了元数据的一致性和优化了并发控制。

主要变化包括：

- 所有提交的元数据现在序列化为 Avro 格式，以便在未来添加新字段时保持兼容；
- 提交的元数据文件名包含完成时间，以支持非阻塞的并发控制；完成的操作、计划和完成的元数据存储存储在基于 LSM 树的、更可扩展的时间线中，使用 Parquet 文件；
- 日志文件格式增加了记录位置信息，支持基于位置的合并；
- 允许表中存在多种基础文件格式，包括 Parquet、Orc 和 HFile，提供了更多灵活性来选择索引和支持新的 ML/AI 格式。
- 此外，1.0 版本还对写入性能、并发控制、二级索引等进行了优化。另外，多表事务的引入使得实现复杂的 ETL 操作成为可能。在实时方面，Hudi 一直处于领先地位，并正在考虑与 Flink 深度结合，实现增量的物化视图能力。

（三）Delta Lake Updates in 2023

在 2023 年的 Data + AI 大会上，Delta Lake 发布了他们的 3.0 版本，这个版本中最引人瞩目的是 Universal Format，简称 UniForm。这一功能允许用户将数据存储为 Iceberg 或 Hudi 的格式，即一次写入可以自动生成多份表的元数据信息。这不仅允许数据存储不必局限于 parquet，还能更无缝地兼容其他数据湖生态。

另一项备受期待的更新是 Liquid Clustering（虽然在 3.0 版本时仍处于预览阶段）。这种新颖的数据组织方式有望提升 Delta Lake 表的查询和管理效率。Liquid Clustering 通过动态调整表中数据的物理布局，优化数据存储模式，以支持高效的数据访问并减少查询延迟。相比之前的静态 order/zorder/hash 等方式，根据查询动态调整可以使用户避免理解复杂的分布式系统原理，从而大幅提升易用性。

此外，Delta Lake 还在架构上进行了优化，抽象出了一个名为 Delta Kernel 的 API。Delta Kernel 旨在简化 Delta Lake 的集成和使用。它提供了一组核心 Java 库 API，使得开发者能够构建和维护 Delta Lake 连接器，而无需深入了解 Delta Lake 的内部协议细节。这使得开发者可以更容易地实现对 Delta Lake 表的读写操作，提高了与不同计算平台的兼容性和灵活性。

从 2.3 和 2.4 版本开始, Delta Lake 对 update 和 delete 进行了大量优化。特别是对 delete vector 的支持,大幅提升了更新和删除的效率,以及提升了 CDC 的易用性。用户可以通过 table_changes 等函数来更容易地获取变更。

在生态集成方面,所有 Delta 连接器(如 Flink、Presto 等)都已集成到 Delta Lake 主仓库,形成统一的生态系统,展现了 Delta Lake 对开放生态的支持。这可能也是 Delta Lake 对 Iceberg 和 Hudi 的追赶。

(四) Apache paimon Updates in 2023

2023 年 3 月 12 日, Flink Table Store 项目顺利通过投票,正式进入 Apache 软件基金会 (ASF) 的孵化器,并更名为 Apache Paimon (incubating)。这标志着 Paimon 从 Flink 的一个子项目,经过一年多的成长,已经演变成一个新的数据湖格式。通过与 Flink 的完美配合, Paimon 逐渐成长为流式数据湖的最佳方案。Apache Paimon 具有强大的 Upsert 更新能力,天然的 DataSkipping 能力,同时还提供完整的 Changelog 流读,配合 Flink SQL 可以实现数据的增量加工。

Apache Paimon 的最新进展包括 0.5 和 0.6 两个版本的更新。0.5 版本增加了数据湖的 CDC 数据摄取成熟度、离线数据仓库的不变视图标签、主键表的生产就绪动态桶模式,以及作为 Hive 表替代的只追加可扩展表。

而 0.6 版本进一步扩展了 Paimon 的特性和功能,包括 Flink Paimon CDC 支持更多主流数据摄取源、跨分区更新、读优化表以提升查询性能,以及新增的 Paimon Spark、Paimon Trino 和 Paimon Presto 模块适用于生产环境,还有与 Flink Metrics 集成的新指标系统。此外, Apache Paimon 加深了与流处理技术的整合,特别是与 Apache Flink 的整合,使得直接在数据湖上进行流处理成为可能,增强了 Flink 的能力并简化了数据处理流程。

(五) 数据湖存储小结

整体来看,数据湖存储还是一个高度竞争力的市场,既有 Iceberg、Hudi 这样的先发者(双方相继在 2021 年末和 22 年初成立了各自的商业化公司),也有 Delta lake 这样由巨头 Databricks 商业化支持的挑战者,以及由阿里支持从 Flink 孵化的流式数据湖格式 Paimon。

在北美和中国,这些数据湖格式的发展状况存在显著差异。

在美国, Iceberg 在生态支持和使用范围上最为广泛,不仅在 Netflix、Airbnb、Apple 等大型企业得到大规模应用并取代了 Hive,还得到了许多云服务提供商和数据

仓库公司的支持（例如 Snowflake 和 Dremio 推出了托管的 Iceberg 表，AWS 和 GCP 也更倾向于支持 Iceberg，而 Azure 更倾向于支持 Delta Lake）。

Delta Lake 凭借 Databricks 提供的完整大数据技术栈成为市场份额最高的商业化产品，并尝试通过 UniForm 来争夺其他数据湖格式的市场份额，而 Liquid Clustering 技术则是 2023 年的技术亮点。

Hudi 在北美相对来说是一个相对小众的选择，但在国内，Hudi 凭借其完整的功能和在实时 upsert 场景下的领先优势获得了字节跳动、快手等大型企业的采用，并且更适应国内对实时场景的强烈需求，在国内取得了良好的发展。

Apache Paimon 吸收了 Hudi 发展过程中的问题（尤其是对实时写入过程中性能和稳定性的复杂索引系统），通过流批一体的方式降低了用户的使用门槛，成为 Hudi 的一个强有力的竞争对手。

四、湖仓计算引擎 2023 发展

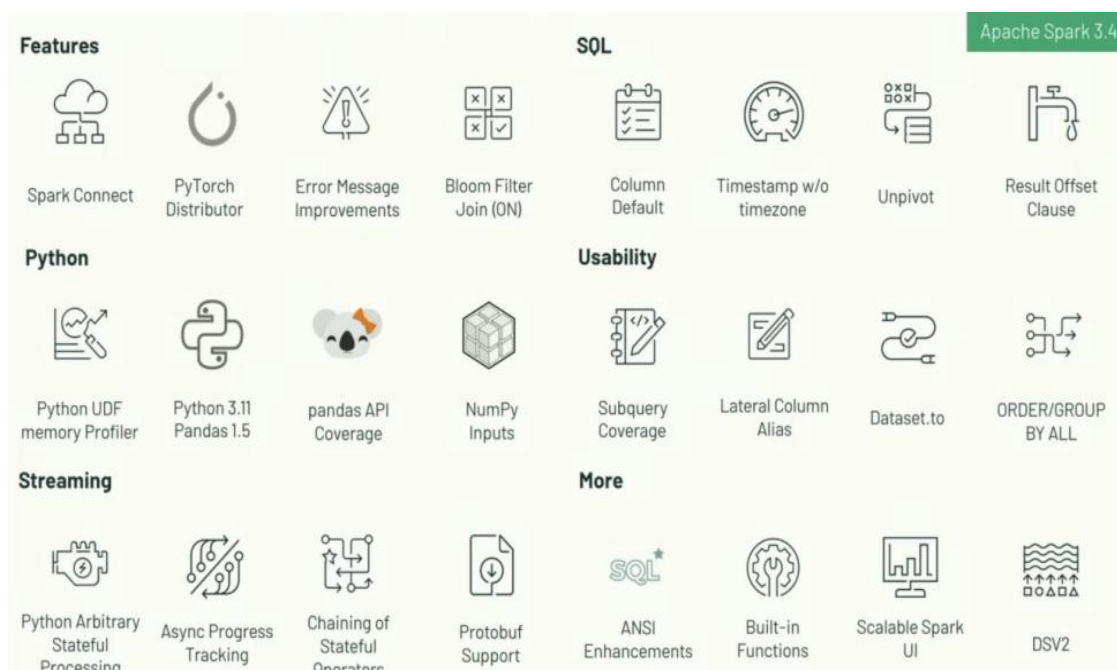
与湖存储中的竞争相比，湖仓的计算引擎层具有更明显的差异化特点，主要可以分为三个赛道：批计算、流计算和交互式分析。当然，它们之间也存在着流批一体和统一分析入口等长期趋势，但目前这些引擎还是相对的泾渭分明、各具特色。

（一）Apache Spark 2023

Spark 应该是整个湖仓架构中最“古老”的组件，从发布到现在已经 10 年，在 2023 发布了 3.4 和 3.5 两个大版本，主要功能有：

- Spark Connect：允许应用从任意位置连接到 Spark 集群，支持 Scala 客户端，分布式训练和推理，Pandas API 在 Spark 上的对等支持，以及结构化流的改进兼容性。
- 分布式训练：在 PySpark 中加入了 TorchDistributor 支持，简化了在 Spark 集群上使用 PyTorch 进行的训练。也支持通过在 Spark 集群上使用 DeepSpeed 简化分布式训练。
- SQL 和 PySpark 新功能，包括多个新 SQL 函数，支持 SQL IDENTIFIER 子句，列支持默认值，新增 UNPIVOT 操作，以及 Python UDFs 的 Arrow 优化。
- 流处理性能与稳定性提升：通过 RocksDB 状态存储提供者和流处理改进，提高了性能和稳定性，并提供了原生对 Protobuf 格式读写的支持。

- 改善 Python 开发者体验：引入新的错误消息框架和内存分析工具，增强了 PySpark 测试 API 和错误类。
- 自然语言 SDK：Apache Spark 的英语 SDK 允许用户使用简单英语作为编程语言进行数据转换，使数据操作更加易于访问和用户友好。



(二) Trino/Presto 2023

Trino 和 PrestoDB 虽然分叉，但整体路线上还是相近的。然而，总结 2023 年的发展情况，这两个公司似乎没有太多突出的进展。首先，它们的目标仍然集中在性能方面。PrestoDB 采用了使用 C++ 的本机引擎的路线，他们测试了使用 Velox 替换 Java 运行时以提高性能约 2 倍，但这个项目已经进行了三年多，但迄今为止还没有合并到主分支中，预计将在明年的 Presto 2.0 中推出。

相对而言，Trino 2023 核心功能上的迭代更快一些：

- Trino Gateway 的推出：作为一个代理和负载均衡器，简化了管理和查询多个 Trino 集群的过程。
- SQL routine：允许用户定义和重用自己的函数，提升了查询的清晰度和简化度，类似 SQL UDF 的功能。
- Schema evolution 和 dynamic catalog：在 Hive 中支持列数据类型的更改、列的重命名等，以及实验性地添加了动态目录功能。

- Project Hummingbird: 做了一系列的性能提升工作。
- Lakehouse 改进: 增加了对 Hudi、Delta Lake 和 Iceberg 的新功能和元数据表支持。
- Java 21 支持: 从 Java 17 升级到 21, 带来了多项改进。
- 在生态系统方面, 主要优化了 Python client 的能力。

(三) Apache Flink 2023

Flink 在流计算领域逐渐成为事实标准, 并在 2023 年发布了两个重要版本: Flink 1.17 和 1.18。这些版本带来了以下主要功能:



- Streaming lakehouse API 对 Delete/update、DDL、存储过程和 Time Travel 的支持。
- SQL client 支持 Gateway 模式, 并支持 JDBC。
- 执行层对执行计划的优化批处理能力提升, 1.18 版本相比 1.16 版本核心性能提升 54%。
- SQL 算子力度的 State TTL 支持。
- Hybrid shuffle 支持远程存储。
- 增量 Checkpoint 成熟可用。
- 实现高效弹性扩缩容新范式, 云原生能力进一步提升。

(四) StarRocks 2023

在 2023 年, StarRocks 实现了从一个单纯的 OLAP 引擎向一个湖仓引擎的转变, 并发布了三个重要版本: 3.0、3.1 和 3.2。

StarRocks 借助存算分离架构的核心, 提升了计算的弹性, 降低了存储成本。

用户可以将 StarRocks 作为一站式 Lakehouse，通过它提供的高速统一查询能力，实现数据导入和查询的一体化。此外，针对开放数据湖（Hive、Hudi、Iceberg、Paimon），StarRocks 可以直接分析数据湖以加速交互式查询分析。如果直接查询数据湖的性能不满足需求，还可以通过物化视图加速查询，而无需进行复杂的数据加工和 ETL 操作。

其重要功能包括：

- 存算分离功能成熟可用，支持 K8s operator 和 Helm 的部署，能够实现秒级弹性。
- 支持存算分离的多 warehouse 的资源隔离，通过不同的仓库来实现不同业务不同负载的物理隔离。
- 支持 Iceberg Hive 的数据湖写入能力，支持 Paimon 的查询能力。
- 支持 Trino 的方言语法兼容，能够更容易地支持 Presto/Trino 的服务切换。
- 支持 Spill 的算子落盘能力，保证大查询的稳定执行，不会出现 OOM，大幅提升跑批的稳定性。
- 强化物化视图能力，支持 Hive/Iceberg/Hudi/Delta lake/Paimon 的外表物化视图能力，通过物化视图的改写能力可以实现透明加速，实现轻量化的 ETL 来降低 pipeline 的复杂度。

（五）计算引擎小结

在数据处理领域，各个计算引擎都在不断发展和创新，以满足不断增长的数据处理需求和应用场景。

在数据处理场景下，Apache Spark 在 AI 场景下的训练和生态兼容性上实现了大幅增强，而且通过加强其处理大规模数据集的能力，进一步巩固了其在数据处理领域的核心地位。Apache Flink 则继续在流计算领域保持其领先地位，同时通过增强其云原生能力和提升批处理性能，使得其更加适应现代云计算环境的要求。此外，Flink 与 Apache Paimon 的结合，进一步优化了流式数据仓库的能力，展现出更加强大的实时数据处理和分析能力。

在查询性能方面，Trino 和 Presto 通过持续优化其执行引擎，显著提高了大数据查询的效率和速度，使得用户能够在更短的时间内获得洞察和结果。而 StarRocks 因早已完成了 C++ 原生的向量化改造，则在湖仓一体化的场景中，通过其物化视图的支持，不仅简化了数据加工和查询的流程，还通过其存算分离的架构设计，显著提升了计算的弹性和资源利用率，进而保证了高性能和高效率的数据处理能力。这种多样化的技

术进步不仅丰富了数据湖上的计算引擎生态，也为企业提供了更多的选择和灵活性，以适应不断变化的数据处理需求和业务场景。

五、Lakehouse 的落地和案例介绍

（一）腾讯微信基于 StarRocks + Iceberg 的湖仓一体实践

1.湖仓一体探索

湖仓一体要解决的是通用大数据计算引擎处理 OLAP 数仓体验割裂、存储多份的问题。过去，无论是在湖上面还是仓上面，都存在着相应的困境。

湖上面面临的主要问题是查询性能太慢，如果需要对 Hadoop 中的数据进行即时查询，则需要先将数据导出到 OLAP 数据库中再进行查询加工，造成了资源的浪费，工序繁琐；而仓上面，存在着通用大数据计算引擎处理 OLAP 实时数仓生态不够友好的问题，例如，Spark 无法直接分析 ClickHouse 中的数据，如果要进行查询分析，则需要先将数据从 ClickHouse 落回 Hive 中。总结来看，过去的架构存在的问题就是：

- 两份存储
- 口径不一致
- 额外的数据导入导出步骤
- 数据分析流程难以标准化

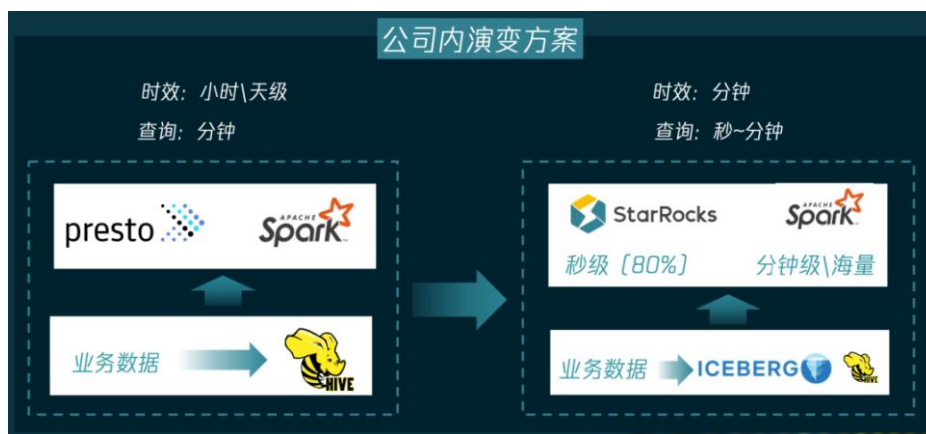
通过湖仓一体架构，微信希望能够真正解决上述问题，实现存储统一、接口统一、元数据统一，亚秒级查询与分钟级查询统一等，最终理想化的湖仓一体形态：面向 SQL，用户不再需要感知底层架构。

2.技术路线一：湖上建仓

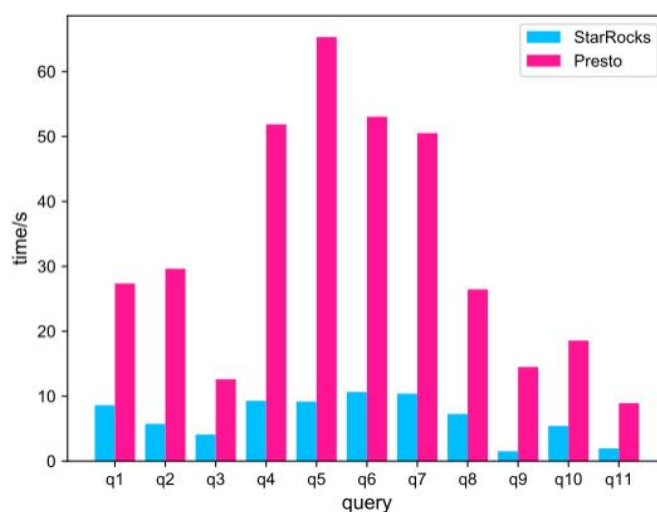
实现湖仓一体的第一种技术路线是湖上建仓，即在数据湖基础上实现数仓的功能，代替传统数仓，构建 Lakehouse。

在微信内部，湖上建仓的架构经历了从 Presto + Hive 到 StarRocks + Iceberg 的演变过程，通过使用 StarRocks 替代 Presto，数据的时效性从小时/天级提高到了分钟级，同时查询效率从分钟级提高到了秒级/分钟级，其中 80%的大查询用 StarRocks 解决，秒级返回，剩下的超大查询通过 Spark 来解决。

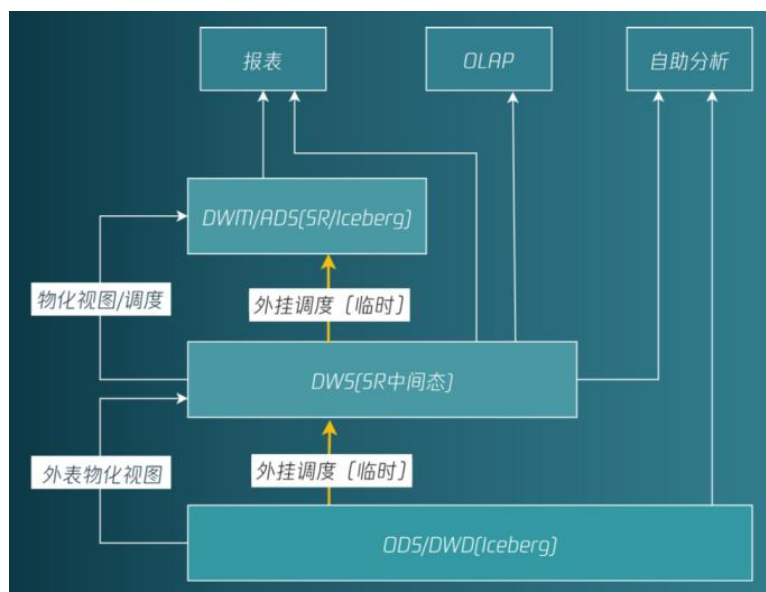
- 秒级：中大表，秒级返回，StarRocks
- 分钟级：大表查询，分钟级，Spark



与 Presto 相比, StarRocks 直接查湖的性能提升 3-6 倍以上。



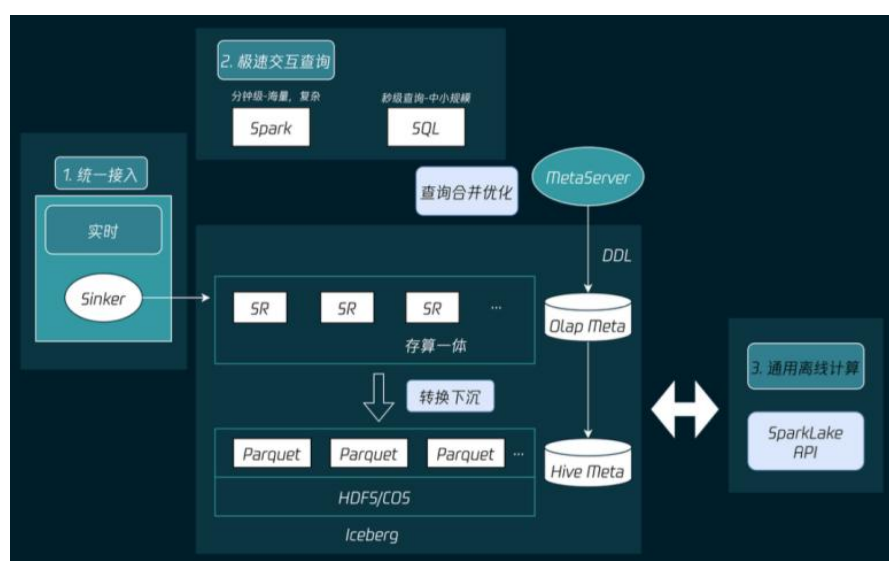
目前, 在微信的使用场景中, 湖上建仓的方案主要通过外挂调度的方式来实现。未来, 当 StarRocks 的外表物化视图功能更加成熟以后, 会采用更加简洁、方便、易于运维的外表物化视图来实现。



湖上建仓方案具有的优点是成本低，实现简单，同时 Hadoop 生态兼容性更好。但其存在的缺点是：数据延迟较高，需要 5-10 分钟左右；另外，ODS、DWD 层的查询会比较慢，需要通过本地缓存等方式来进行加速。因此，综合来看，湖上建仓的方案主要还是湖为主，更适合于离线分析为主的场景。

3.技术路线二：仓湖融合

实现湖仓一体的第二种技术路线是仓湖融合，通过在数仓中加入跨源融合联邦查询的功能，打通内容存储，从而不需要经过 ETL 能够直接分析数据湖。在该方案中，微信采用冷存下沉的手段，实现仓湖数据的统一。



在该方案中，数据会先实时接入仓，之后，冷存经过转换下沉到湖，通过 Meta Server 实现仓湖元数据的统一管理，在查询时，能够合并冷热数据读取，同时，通过 SparkLake API 提供对通用离线计算的支持。

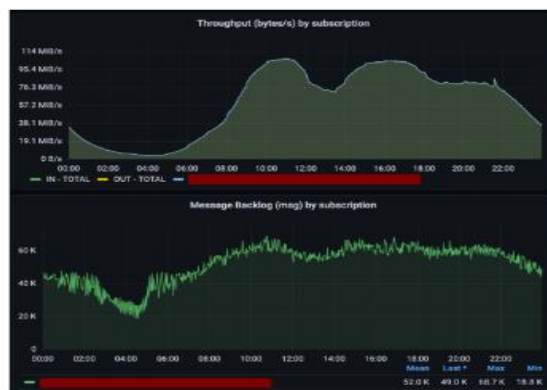
在该方案中，由于数据是先落仓的，因此数据的实时性会更高，在秒级到 2 分钟以内，同时 DWD 层的查询响应会更快。但其缺点是成本高（热数据 TTL），Hadoop 生态兼容性不如湖上建仓的方案。

因此，综合来看，仓湖融合的方案主要还是考虑仓为主，更适合于实时分析为主的场景。

当前，仓湖融合、冷存下沉的湖仓一体方案已经在微信安全的业务场景中落地，接入的大表每日数据量达数十亿，超大表每日数据量千亿以上。小时分区降冷耗时分钟级，天分区降冷耗时小时级；在内存消耗方面，单任务最大消耗 5GB 左右，对集群无明显影响。同时，通过 BE/FE 参数调优：compaction 线程数比例调整、flush 线程数调整、并发事务数调整，消除了大表并发时的数据挤压，全天稳定运行：



上线前



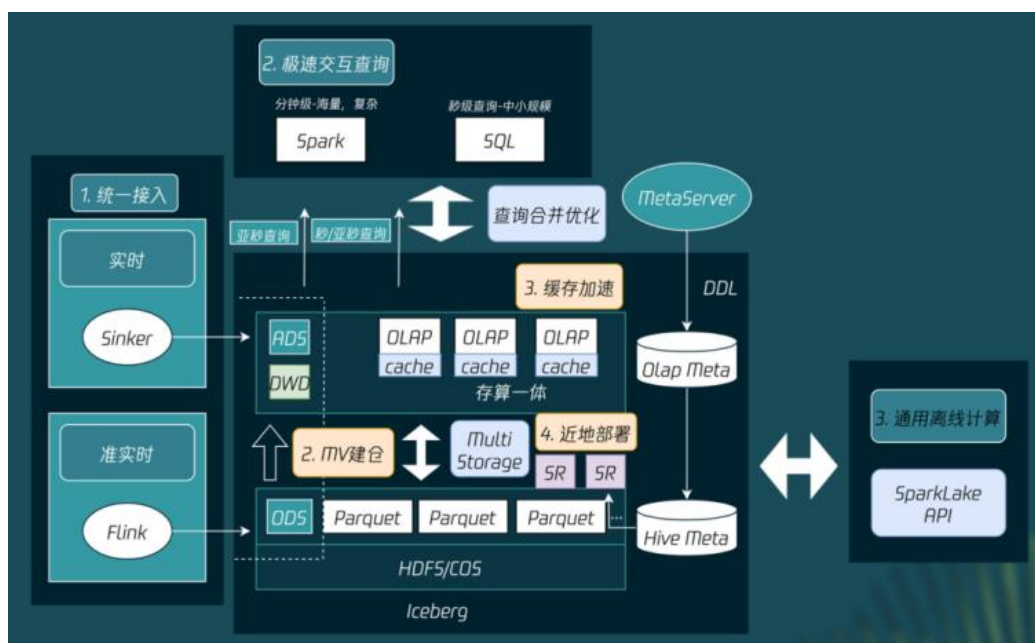
上线后

4.湖仓一体架构

从前面的分析中可以看到，湖上建仓和仓湖融合两种方案各有优缺点，分别适用于湖为主和仓为主要的业务场景，而在微信的业务场景中，两种路线都有需求。

方案	湖上建仓	仓湖融合
数据时效性	5~10 分钟	秒~2 分钟
DWD 查询响应	秒	亚秒
ADS 查询响应	亚秒	亚秒
Hadoop 生态兼容性	优	一般

因此，在微信的湖仓一体方案中，采用的是湖上建仓和仓下沉湖的融合方案。



在该方案中，微信同时支持实时接入和准实时接入，用户可以根据自身对成本、性能和数据时效性的要求来选择通过哪种方式进行接入，并且是可切换的。例如，刚开始时，用户对性能、数据时效性要求不是特别高，那么可以选择通过湖接入的方式，这样成本较低，当后续该用户对性能和数据时效性要求变高了，那么他可以切换到仓接入的方式，这样就可以满足需求。因此，在微信的湖仓一体架构中，既能够支持极速 BI 分析的场景，也能够满足通用离线计算的需求。

当前，基于 StarRocks 的湖仓一体方案已经在微信的多个业务场景中上线使用，包括视频号直播、微信键盘、微信读书和公众号等，集群规模达到数百台机器，数据接入量近千亿。以微信一个具体的业务举例，在直播业务场景中，通过湖上建仓的方案改造，使得数据开发同学需要运维的任务数减半，同时存储成本降低 65% 以上，离线任务产出时间缩短两小时。

（二）Lakehouse 的架构建议

基于当前的存储引擎和计算引擎的发展状态，如果在 2024 年选择湖仓架构，以下是两个推荐选项：

1.Iceberg + StarRocks 方案：这是一个理想的选择，特别适用于需要高性能查询和灵活数据管理能力的场景。Apache Iceberg 作为一种开源的表格格式，具有清晰的抽象层次和出色的扩展性，已经成为全球范围内使用最广泛的湖格式之一。Iceberg 不仅可以有效替代 Hive，还能提供更高效的数据存储解决方案。而 StarRocks 作为分析引擎，提供比 Trino 更出色的性能和资源隔离能力。通过物化视图，StarRocks 可以大大简化开发工作量，并且内置的主键模型可以实现高效的实时数据更新。此外，StarRocks 还具有向 Iceberg 反向降温数据的能力，可以有效补充 Iceberg 在实时处理能力方面的不足，这也是一些大型企业选择该方案的重要原因。

2.Flink + Paimon 方案：这个方案更适合需要强大实时处理能力的场景。Apache Flink 已经在实时计算领域确立了自己的地位，而 Paimon 作为一种新兴的存储格式，具有出色的性能和良好的国内社区支持，成为与 Flink 结合使用的理想选择。这种组合不仅可以提供紧密的集成和优异的实时数据处理能力，还为未来流式数据湖的建设留下了足够的空间。

总结：

在 2010 年代，数据湖的概念应运而生，其目标是打破数据孤岛，解决由封闭的多个数据集市和数据仓库带来的数据管理难题。随着 Hadoop 的开源生态的发展，数据湖的理念逐渐得到实践和验证。

进入 2020 年代,随着底层湖存储能力的提升,例如 Iceberg、Hudi、Delta Lake 和 Paimon 等,以及湖仓计算引擎如 Spark、Trino 和 StarRocks 的发展,数据湖开始逐步替代传统的数据仓库,实现了湖仓一体化。这种新的架构模式不仅提高了数据处理和分析的效率,也使数据的管理和维护更为灵活和方便。

在未来,湖仓一体化的发展趋势将更加明显。随着大数据技术的进一步发展,我们期待看到更多的创新和变革。例如,实时数据仓库的发展将使得数据处理和分析更加实时和精准,为企业提供更 valuable 的洞察。同时,随着 AI 和机器学习技术的发展,湖仓一体化架构可能会引入更多的智能化元素,提高数据处理的自动化程度。

16、数据库兼容性

一、背景

随着国产数据库逐步成熟及信创动作的扩展下沉,越来越多企业计划或已开展国产数据库替代工作。根据墨天轮 2023 年发布的调研数据显示,有超 86%的企业有替换国外数据库的计划,其中有近 60%正在使用 Oracle 数据库的企业想替换成国产数据库。

Oracle 作为市场占比最高的数据库系统,无疑是国产数据库的标杆,历经数十年的发展,Oracle 已建立一套成熟且卓越的技术体系,如何平滑、安全、低成本地完成从 Oracle 到国产数据库的迁移工作,是国产数据库厂商必须解决的问题,其中首要解决的就是与 Oracle 的兼容性问题。

Oracle 作为成熟的数据库管理系统,具有丰富的功能和强大的性能。国产数据库要实现与 Oracle 的兼容,需要解决一系列技术难题,如数据类型、SQL 语法、存储过程等。因此 Oracle 兼容性不仅是模仿,而是一个非常复杂和工程量庞大的逆向工程。以 Oracle 19c 为例,在 SQL 层主要功能罗列如下:

1. 语法和常见功能函数;
2. 结构化数据类型;
3. JSON、XML 等半结构化数据及功能;
4. 查询加速提升,如 OLAP、并行处理、结果集缓存等;
5. 如 UDT、高级包、存储过程、触发器等各种高级特性;
6. 安全、加密和审计;
7. 工具生态上大量的衍生工具和中间件等。

对于国产数据库厂商而言，在 Oracle 兼容方面需保持务实的态度，将重点放在核心功能的实现上，并保持高度的兼容性与稳定性。

二、实践

数据库是一个极其复杂的软件工程产品，SQL 标准仍在不断发展，诸如 ISO/IEC 9075:2023 等新标准不断涌现。因此，从技术演进角度，直接复用已有的组件并不利于数据库产品的持续演进。

所幸的是，随着国产数据库逐步成熟，无论是传统数据库厂商，还是新型数据库厂商都表现出更多的自信，以“兼”攻“替”。尤其是新兴的全自研国产数据库厂商，自主研发打造 SQL 引擎，体系化兼容 Oracle，实现语法、语义、PL/SQL 高级包以及工具生态四个层面的高度兼容。

（一）自主 SQL 引擎助力全面兼容

SQL 的本质在于理解用户编写 SQL 语句的意图，然后转成高效执行计划，然后操作和管理数据库执行要求的运算算子，返回对应结果集。其过程需要实现编译、优化、执行 3 部步曲。需要理解 SQL 文本的意图，然后转成高效执行计划，再进行数据库的高效执行算子。

对于 YashanDB 的 SQL 引擎来说，为了快速实现 Oracle 兼容性，从架构实现上采取了有效的方法：通过对 SQL 语法、函数、数据类型、存储过程语法、高级包、优化规则、执行算子等各个维度进行抽象和归纳，在 SQL 引擎的各个功能阶段提供功能注册点。由于全链路的代码都是自主研发，可以在框架内通过扩展功能注册点来快速添加各种兼容性功能，对于部分涉及调整框架的需求，也能够确保整个架构的稳定性。

（二）Oracle 兼容性原则

数据库的高级特性往往涉及较高的技术挑战，如 UDF(用户自定义函数)、JAVA UDF、C UDF、存储过程、高级包、UDT、触发器、JOB、DBLINK 等。

1.PL/SQL 语言的功能

这是 Oracle 绑定很多用户应用的锁链，为了降低用户应用迁移成本，YashanDB 实现了大部分的 PL/SQL 语句，如声明、赋值、控制、跳转、循环、游标处理、SQL 语句、动态 SQL 以及异常处理等。同时提供了全面的过程体对象，如存储过程、UDP、UDT、触发器等供用户使用。以 UDT 自定义类型为例，YashanDB 提供了 RECORD、OBJECT、VARRAY、NESTED TABLE 多种形态，支持方法的声明、支持将嵌套表形态。

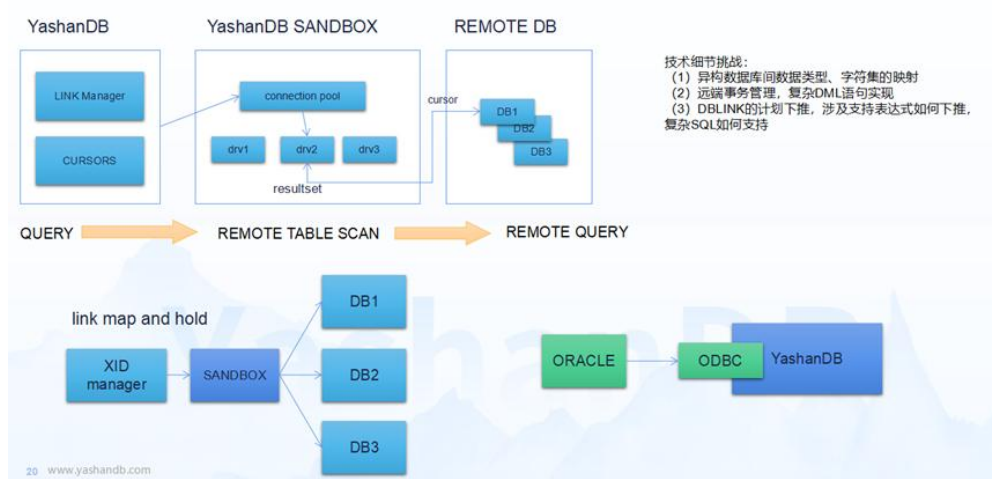
2.JAVA UDF 方案

Oracle 为了实现这一特性，直接将 JVM 集成到内核中；由于 Oracle 采用多进程架构，一旦用户编写的 UDF 存在异常，影响范围有限，可以重新启动进程。然而，由于 YashanDB 采用多线程架构，一旦用户要运行的代码出现异常，对我们来说是毁灭性的冲击。所以我们需要设计一个沙箱进程，由该进程负责拉起 JVM 执行用户 UDF。

3.DBLINK 功能

YashanDB 提供了从 YashanDB 到 YashanDB，YashanDB 到 Oracle，Oracle 到 YashanDB 的远端访问和事务能力，在这个功能点上，存在很多技术细节需要攻克，如异构数据库间数据类型、字符集的映射，计划和表达式如何进行远端下推，复杂 SQL 的拆分，远端事务管理，复杂 DML 语句如何联动等。

YashanDB Oracle兼容性DBLINK



除了 SQL 功能角度，YashanDB 还从以下方面完善 Oracle 兼容性要求，包含统计信息，收集基于列的统计和直方图信息；计划优化规则；常用字符集支持，UTF8、GBK、UTF16、18030 等常见字符集支持；安全、审计与加密；导入导出工具支持；JDBC/ODBC 驱动等。

总结：

在过去的几十年时间里，Oracle 数据库凭借出色的产品能力与硬核的技术实力始终处于领先地位，使得国内诸多关键行业的企业级应用深度依赖 Oracle。数据库和应用的不兼容，会导致用户需要付出极大的迁移成本，带来的直接后果就是难以形成规模化复制。在实现 Oracle 兼容性的目标上，YashanDB 总结出以下方法供同行参考：

- (1) 特性对标，寻求适合技术架构的落地方案；
- (2) 特性细节，从功能特性角度，摸清实现细节，从细节中提炼特性的实现原则；

- (3) 兼容不是抄袭、不能盲从，甄别 Oracle 的特性规格和特性实现 BUG；在开发过程不断汇总实现细节上的典型案例，同时编写开发指导文档提供用户避免踩坑；
- (4) 预留扩展路线，为后续技术竞争力落地做好铺垫。

17、数据库内核创新

一、数据库内核创新：高度可扩展内核架构

（一）概述

当前，人类正在经历新一轮的由 AI 引领的信息革命。新一轮的信息革命为人类提供了新的生产力工具，有望带来新一轮的生产力大发展。在这新一轮的信息革命中，数据依然扮演着最为重要的角色。数据库作为存储与处理数据的关键技术，在这新一轮的技术浪潮中，也越来越凸显其价值和重要性。这也对数据库提出了更高的要求。

1. 数据类型比以往更多样。组织内不再只是有表格数据（也就是我们所熟知的关系型数据），还有大量的文档、图片等非结构化的数据，如何更加高效、便捷地处理这类数据也是现有数据库需要重点考虑的问题；
2. 新技术带来的新的需求，如 AI，尤其是 LLMs（大语言模型）的发展，对向量数据的处理有了更高的要求，现有数据库技术如何升级来满足此类需求？
3. 工业自动化比如智能制造的发展，也使对现有数据类型的处理有了更加细分的需要，比如时序数据。时序数据本质上来说也是关系型数据，但由于其数据来源主要是来自设备，因而有数据量大、采集频率高、数据质量不稳定等特点，因而也需要使用相应的数据处理技术来应对此类需求；
4. 组织面对的业务场景更加多样，数据连接的维度要求更高了，分析此类问题，使用图技术是个很好的方式。现有数据库技术如何融合图技术也是个值得思考的问题；
5. OLTP 和 OLAP 的边界越来越模糊，一个数据库往往既要承担 TP 的任务，同时也要承担 AP 的任务，即所谓的 HTAP（Hybrid Transaction/Analytical Processing，混合事务分析处理）。如何打造一个合格的 HTAP 型的数据库是每个数据库厂商的必修课；
6. 对数据安全也有了更高的要求，隐私也越来越得到更多人的重视，数据库如何实现隐私计算也会是不久的将来不得不考虑的问题；

7. 由于组织内需要处理的数据类型越来越多样，为处理多样的数据而引入了过多的数据库产品种类，但这也带来高昂的学习和管理成本，也使得系统架构变得复杂，增加了安全风险，加剧了“数据孤岛”的形成；
8. 边缘数据计算的需求日益强烈，如何在有限的算力下实现高效、准确的数据处理也是当前数据库需要着重考虑的问题之一；
9. 对数据库自身的，比如故障自愈、系统性能优化也有了更高的要求，下一代数据库的运行不再需要专业的运维管理人员，通过系统自身的能力便能很好地稳定运行。

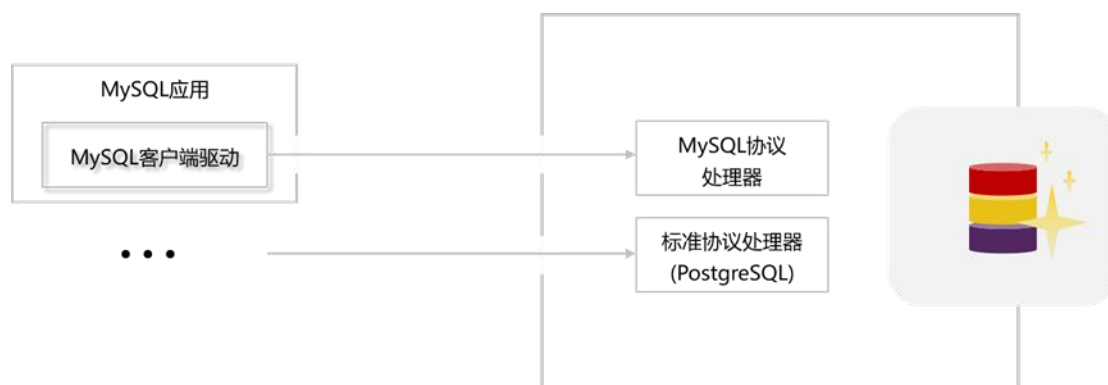
综上所述，数据库需要面对的数据处理场景已经变得更加多样、更加复杂、要求也更高了，所以如何使数据库能高效地去适配这么多场景和需求就显得异常重要了。要达到该目标，其中一个途径就是要设计一个高度可扩展的数据库内核。

（二）现状

数据库是一个非常复杂的软件系统。数据库的核心功能是存储和管理数据，为用户提供高效、安全、可靠的数据处理服务。而随着数据量的增加，数据种类的增加，数据处理需求的多样化、要求更高的数据处理时效以及系统管理的简单化甚至无人化，把对数据库的能力要求推向了一个新的高度。从全球范围来看，目前该领域的佼佼者仍然是美国的 Oracle。从公开的资料来看，Oracle 是一款真正融合了关系型数据，图片、文档等非结构化数据，JSON、XML 等半结构化数据，也拥有图处理能力，并提供一定的故障自愈、系统性能优化能力的 HTAP 数据库。由于 Oracle 是闭源的商业数据库，其内核如何实现我们无从得知，但可以推断的是 Oracle 的内核必定是一个高度可扩展的架构设计。开源数据库领域，PostgreSQL 的 Extension 机制是比较典型的并且也是比较优秀的可扩展内核架构设计，但由于其设计颗粒度较粗，因此也限制了其扩展的能力。

（三）架构原理及实现

要实现高度可扩展的内核，其中的一个重要思想是可插拔（Pluggable），就像计算机上的 USB 接口一样，可以接入完全不同的设备（只要遵循 USB 接口标准）。这就要求对功能模块进行抽象并标准化，形成标准的功能模块接口。只要遵循并实现功能模块接口，就能够实现内核能力的扩展。



以数据库协议模块为例，模块完成诸如连接建立、数据发送、数据接收等基本的、通用的功能，而将数据解包、数据封包等操作抽象成接口。不同的协议根据自己的协议规则实现自己的数据解包、数据封包等操作，从而使得在一个数据库内能支持比如 MySQL、PostgreSQL 等多个协议。

对高度可扩展内核的探索，目前国内外都比较积极。比较典型的如 AWS Babelfish 实现了对 SQL Server TDS 协议的支持。国内的羲和（Halo）是国内首个支持多协议（MySQL/PostgreSQL）的数据库产品，并且也是国内首个在查询编译器、查询优化器、查询执行器等层面实现 Pluggable 的数据库产品。

（四）趋势及挑战

实现高度可扩展内核是数据库技术发展的必然趋势，这是不断增长的数据处理需求决定的。不过实际应用过程中，仍然面临诸多挑战：

- 性能优化。数据库是一个性能敏感的系统。可扩展性的实现可能会引入额外的开销，如上下文切换、内存分配等。为了保证系统的性能，研发人员需要发扬工匠精神，下功夫在某些技术细节层面深入挖掘。
- 系统稳定性。可插拔的内核模块可能会影响系统的稳定性。因此，在设计内核时，需要充分考虑模块间的依赖关系和资源管理问题。

未来已来，面对纷繁复杂的数据处理场景，当前的数据库也需要一次自我进化。下一代数据库不再是单纯的数据库管理系统了，而应该是融合的、集 AI 能力为一体的数据操作系统（Data Operating System, DTOS）。

二、数据库内核的微观评测与优化

（一）什么是数据库内核的微观评测

数据库内核的微观评测，不同于目前现有的数据库评测体系。现有数据库评测体系，无论国际非营利组织 Transaction Processing Performance Council (TPC，事务处

理性能委员会)的 TPC-C 压测，还是各种开源或商业评测工具，如 Sysbench、Hammerdb、PostgreSQL 自带的 pgbench。无一例外，都是采用执行某些特定类型的 SQL、存储过程，同时记录每秒的查询数（QPS）、事务数（TPS），并以最终的 QPS、TPS 为测试结果。这种以 QPS、TPS 为准的，都属于宏观评测。

微观评测不以整条 SQL 为评测对象，而是将一条 SQL 的执行，分成若干细小的阶段，如 SQL 语句解析阶段、逻辑读取阶段等等，甚至细分到内核中一个仅有几百条指令的函数，都可以成为评测对象。微观评测也不以 QPS/TPS 为衡量标准，取而代之，换为如下指标：指令数、周期数、L1 命中率、STALL 周数、分枝预测命中率、预测失误周期数等单一指标，或是 IPC（每周期指令数）、CPI（每指令周期数）、Mem Stall 周期、L1 Miss Stall 周期等复合指标。

（二）数据库宏观评测与微观评测的比较

宏观评测的优势是简单、明了。绝大多数的宏观评测，都十分简单。Sysbench、pgbench，都是一条命令就可以开启压测。其压测结果就是 QPS/TPS，十分容易理解。但在性能分析领域，有一条金科玉律，间隔过久的性能指标，由于“尖峰”、“低谷”会被平均，问题将被掩盖。

宏观评测最基础的观测对象是 SQL，一整条 SQL。微观评测的观测对象可以细到只有几百条指令的一小段重要内核内码。一整条 SQL 对于成熟的数据库（Oracle/PG/MySQL），最少情况下需要 10 万条量级的 CPU 指令，相比微观评测的百“条指令量级的观测对象，宏观评测的观察对象是微观评测的 1000 倍。这也是”微观“二字的来源。

由于微观评测的观察对象足够的小，被隐藏的“尖峰”、“低谷”，将被充分暴露。这就好像用电子显微镜观察光滑的玻璃镜面，能看到其表面凹凸不平。但在宏观领域，我们凭眼去看，只能看到玻璃镜面光滑无瑕。

宏观评测虽简单，但简单带来的后果，就是评测结果易被操纵。因此，就算是国际非营利组织 TPC 的 TPC-C 压测，其评测结果认可度也并不强。世界最大商用数据库公司 Oracle，在 2010 年之后不久，就放弃了 TPC-C 压测。在 2020 年前后，先后有数家公司重启 TPC-C 打榜之路，虽已经多次取得世界第一头衔，但几年下来，这几个数据库营收累计之和，也就刚刚和 Oracle 一个月的营收持平。这从侧面上反映了 TPC-C 的“世界第一”已无意义。反而是许多小巧的开源或商业宏观压测工具，以简单的使用，简单、易理解的输出结果，被广泛使用。

宏观评测，是针对普通用户的。使其可以快速了解某一个数据库基本的处理能力。结果虽不严谨，但必不可少。就像显微镜能明察秋毫地看到更细微东西，但我们不能随

时随刻抱着显微镜看一切东西。微观评测就如同显微镜，它的评测结果更加专业，多数结果不适合普通用户，只针对数据库内核开发人员。

（三）为什么微观评测是必须的

虽然宏观评测可以使用户、DBA 快速了解一个数据库的能力范围。但因“宏观”之大，而必然产生的疏漏，使我们无法真正了解数据库的处理能力。比如，业内常常拿来对比的 MySQL、PG。有无数人、无数次地做过 MySQL、PG 的横向压测对比。但无论是谁，都无法说清二者到底孰强孰弱。

从宏观层面，二者各有优缺点。你用肉眼去看两块玻璃镜面，是无法分清谁比谁更光滑的。但由于制造时的工艺、温度、湿度等因素影响，二者一定有高下之分。若从微观层面、放在电子显微镜下观察，就能区分出高下。MySQL 和 PG 在事务开始时，都有一个获得 XID(事务 ID)的过程。对应代码所用 CPU 指令数 1000 条出头。在 MySQL 中，对应代码如下：

```
3037 void trx_set_rw_mode(trx_t *trx) /*!< in/out: transaction that is RW */
3038 {
3039     ut_ad(trx->rsegs.m_redo.rseg == 0);
3040     ut_ad(!trx->in_rw_trx_list);
3041     ut_ad(!trx->is_autocommit_non_locking(trx));
3042     ut_ad(!trx->read_only);
3043
3044     if (srv_force_recovery >= SRV_FORCE_NO_TRX_UNDO) { ...
3045
3046     /* Function is promoting existing trx from ro mode to rw mode. ...
3047
3048     trx_assign_rseg_durable(trx); //分配回UNDO段
3049
3050     ut_ad(trx->rsegs.m_redo.rseg != 0);
3051
3052     mutex_enter(&trx_sys->mutex);
3053
3054     ut_ad(trx->id == 0);
3055     trx->id = trx_sys->get_new_trx_id(); //return (trx_sys->max_trx_id+);
3056
3057     trx_sys->rw_trx_ids.push_back(trx->id);
3058
3059     trx_sys->rw_trx_set.insert(TrxTrack(trx->id, trx));
3060
3061     /* So that we can see our own changes. */
3062     if (MVCC::is_view_active(trx->read_view)) { //视图
3063         MVCC::set_view_creator_trx_id(trx->read_view, trx->id);
3064     }
3065
3066     UT_LIST_ADD_FIRST(trx_sys->rw_trx_list, trx);
3067
3068     ut_d(trx->in_rw_trx_list = true);
3069
3070     mutex_exit(&trx_sys->mutex);
3071 }
3072
```

观测目标，共 1137 条指令

图 1

PG 中代码不再列出，读写事务开始时获得 XID 的代码，可自行在 PG 中查找。使用内核微观评测方法，对二者进行分析，结果如下：

	A	B	C	D	E	F	G	H
1				MySQL			PostgreSQL	
2	Instructions			1137			1251	
3	cycle			12017	10.57		9096	7.27
4	mem read number			291	0.26		316	0.25
5	mem write number			249	0.22		262	0.21
6	ITLB_MISSES Cycle			1593	1.4		1383	1.11
7	DTLB_LOAD_MISSES Cycle			2864	2.52		1313	1.05
8	DTLB_STORE_MISSES Cycle			564	0.5		218	0.17
9	Execute Stall			1360	1.2		1056	0.84
10	Mem Stall			6119	5.38		3894	3.11
11	Totale Stall			10027	8.82		7494	5.99

图 2

第一行是 CPU 指令数。二者使用指令数相差无几，一个是 1137 条指令，一个是 1251 条指令。第二行是所用周期，和 CPI（每指令平均所耗周期），MySQL CPI 为 10.57，而 PG 只有 7.27。PG 使用了更多的指令数，但所耗周期数更少。MySQL 使用更少的指令，但所用周期数更多。从微观上观察，PG 胜出。但是什么原因导致 PG 胜出呢？

可以继续使用内核的微观评测方法进行观察、分析。第三、四行，是内存读、写的次数，和指令数与内存读写次数的比值。MySQL 一条指令对应 0.26 次内存读，也就是平均 4 条指令对应一次内存读。PG 和 MySQL 差不多，一条指令对应 0.25 次内存读。内存写方面，MySQL 与 PG 也相差不多。

区别最大的是第 5、6 行，DTLB_LOAD_MISSES Cycle 和 DTLB_STORE_MISSES Cycle，第 5 行，读内存时 TLB Miss 产生的延迟，MySQL 每条指令对应 2.52 个周期的 TLB Load Miss 周期，而 PG 只有 1.05 周期。MySQL 是 PG 的两倍多。第 6 行是 TLB Store Miss 产生的延迟，MySQL 每条指令对应 0.5 周期的延迟。而 PG 每条指令只对应 0.17 周期的延迟。第 5、6 行是复合指标，主要对应局部变量、全局变量读、写，产生的 CPU 内部 Stall（停顿）。

MySQL 每读、写一次局部变量或全局变量，产生的 STALL 是 PG 的数倍以上，这是 PG 使用更多的指令、但周期数更少的主要原因。熟知 C 与 C++ 编译器的，可以得知，这主要是 C++ 的 Class、继承等机制，会悄悄地让变量分布得更散。而读、写更为分散的变量，必然导致 TLB Miss 次数大大增加，由此，使 MySQL 平均每条指令对应的访存延迟远高于 PG。这其实是语言，而并非数据库本身所造成的差异。但无论如何，这都是数据库的差异。

下面再使用数据库内核微观评测方法，以另一段高频代码为对象进行评测。评测目标是 MySQL 和 PG 在 Buffer Pool 或 Shared Buffer 中搜索 root 块地址的代码。

I	J	K	L	M	N	O	P
			MySQL			PostgreSQL	
Instructions			538			694	
cycle			6772	12.59		5709	8.23
mem read number			160	0.3		167	0.24
mem write number			107	0.2		118	0.17
DTLB_LOAD_MISSES Number			12	0.02		14	0.02
DTLB_LOAD_MISSES Cycle			1772	3.29		1359	1.96
Mem Stall			4029	7.49		3520	5.07
L1 Cache Hit Pct			89.30%			83.80%	
L2 Cache Hit Pct			89.30%			85%	
L3 Cache Hit Pct			93.80%			93.40%	

图 3

通过第二行的周期数、CPI，可知结果仍和刚才一样，PG 使用更多的指令，但周期数更少。这里对 L1/L2/L3 的缓存命中率也进行了观测，两个数据库类似。TLB Load Miss 导致的 Stall 周期，仍是二者差距的产生原因。对更多高频代码段进行分析，结论基本一致。可知微观层面，PG 的代码质量，强于 MySQL。

但宏观层面，由于二者算法不同，微观层面的差异被掩盖。比如 MySQL 使用了“索引组织表”（InnoDB 引擎以主键索引顺序存储数据），而 PG 的堆表是无序的。使用主键索引搜索，即使 PG 在访存效率上更好，但 MySQL 所用指令数要明显少于 PG，所以性能还是超过 PG 的。而如果使用非主键索引搜索，MySQL 又慢于 PG，主要原因仍不是 PG 访存效率更高，而是非主键索引搜索，MySQL 使用的指令数又多于 PG。

但是，如果 PG 中的行被 Update 过几次，产生了“行链”，如果行链比较长，甚至跨块了，那么即使非主键索引搜索，PG 的指令数又会多于 MySQL，性能也会慢于 MySQL。

这些差别都是数据库特性导致的，数据库内核的微观评测，就是为了排除掉这些特性造成的差异，让我们看到不同数据库最本质的差异。

（四）微观评测的意义与作用

计算机体系结构领域，有一本经典教材：《计算机体系结构量化研究方法》。由图灵奖获得者撰写，介绍体系结构的评测方法。计算机体系结构要图灵者获得者告诉我们如何评测，同样是基础模块的数据库，且离应用最近，往往被称为基石的数据库，远不是随便装一个 sysbench 等软件，跑若干时长的压测，就能得出谁好谁坏？

数据库的评测体系目前尚不完善，难以分辨优劣，不单是用户难以选择，也导致个别国产数据库浑水摸鱼。采用数据库内核的微观评测方法，对高频代码段进行分析，可以最大程度地分辨李逵和李鬼。

除此之外，内核的微观评测方法，相当于对内核高频代码段进行了详细的 Profiling，对于数据库内核开发人员，能详细地指出瓶颈所在、是否有调优的可能。这对于增强内核、开发更为强大的数据库，有着重要意义。

18、数据原生时代的数据库

一、什么是数据原生（Digital Natives）

数据原生（Digital Natives）是 Marc Prensky 在 2001 年的文章 “Digital Natives, Digital Immigrants.” 中提到的概念，和数字移民（Digital Immigrants）同时出现。数字原生又可以译作数字原住民或者数据土著，最早在 Marc 的文章中是用于定义在数字化时代长大的人，意为 80 后甚至再年轻些的这代人，一出生就面临着一个无所不在的网络世界，对于他们而言，网络就是他们的生活，数字化生存是他们从小就开始的生存方式。而数据移民则是长大后才接触数字产品并有一定程度上无法流畅使用的族群，他们在诸如游戏、短视频、社交媒体的使用过程中是学习者。

在当今社会，数据原生和数据移民的概念发生了较大的变化，如在信通院云大所所长何宝宏的 2023 年新书《数据原生》中，数据原生和数据移民的概念从“人”引申到了“组织（企业）”层面，即数字原生是指具有数字技术、数字思维和数字经济的特征和能力，能够适应和驾驭数字化时代的个体或组织。数字移民则是指正处于数字化转型中的个体或组织，数字原生是数据移民这类企业数字化转型的愿景和目标。

二、数字原生时代的来临

目前中国除了企业面临数字化转型的诉求外，各行各业的部门和组织也需要对数据高效利用。基于此，本文对数字原生的定义做出了更宽的扩充，从企业与组织扩展到了行业机构。

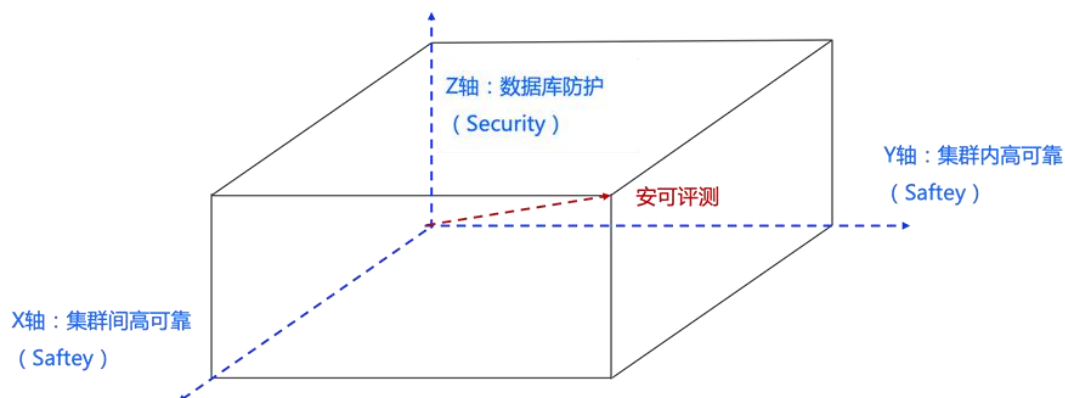
下图对国内较为熟知的 15 个行业按“数据亲和性”进行了划分和定义。（国内除了存在数据移民的企业与组织外，还存在一部分对数据极度不敏感或者尚未引入信息化手段的行业机构/企业，本文引入“数据新民”的概念来描述。）

名称	数据亲和	定义	举例	来源参考
数据原生	高	天然对数据敏感且主要工作是依赖数据的行业机构组织。	金融、人社、统计、应急、气象等。	《“十四五”规划纲要》中提到“优先推动企业登记监管、卫生、交通、气象等高质量数据集向社会开放”。
数据移民	中	以传统实业为主导，对数据较为敏感且在生产过程中产生了大量数据，迫切需要利用这些数据反过来对业务本身进行提能增效的行业机构组织。	能源、医疗、交通、水利、制造、教育等。	《“十四五”规划纲要》中提到“加快推动数字产业化”的方面，如智慧能源、智慧政务、智慧教育、智慧医疗、智慧制造等。
数据新民	低	以传统实业为主导，对数据不敏感，数据和业务系统还处于未治理状态的行业机构组织。	农牧、建筑、餐饮等。	亿欧智库《2022 中国建筑行业数字化转型研究报告》发现建筑行业的数字化程度较低为6.5%和餐饮行业持平，仅高于农牧业 1.9%。

三、数字原生时代的数据库

数字化时代，行业由于对数据的亲和性不同，对数据库产品的能力要求也不同，例如数据原生组织已经完成了数据的集约化，更多考虑的是如果高效将数据共享出去。而数据移民组织具有传统的“本职工作”，数据是工作中产生的附加属性，他们需要将这部分数据集约、收拢，为原业务赋能，实现降本增效的目的。

另外，信息系统的供应链安全已经成为新时代的必要项，随着财政部《数据库采购需求标准》发布，安全能力已经成为行业的标配。以下针对客户提到的“安全”进行了AFK 方法分析扩展：



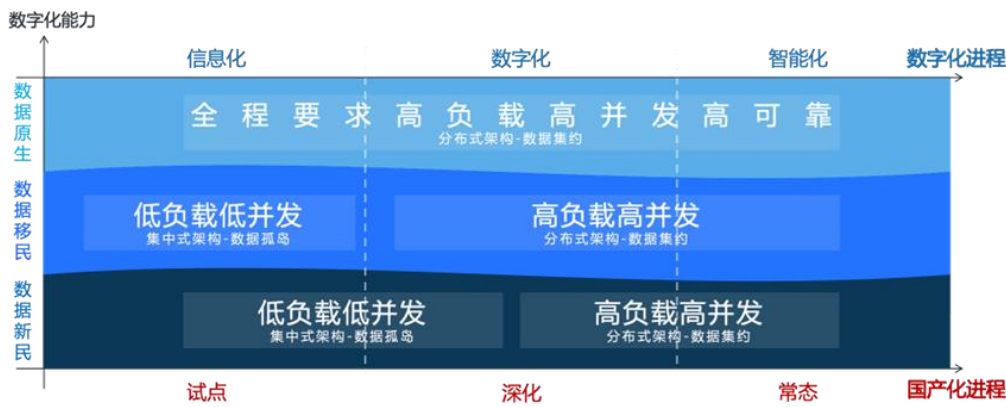
对数据库的安全主要分为两个维度-“大安全”和“小安全”。此两者分别是供应链的安全以及数据本身的安全。数据本身的安全又分为 Safety 和 Security，前者是对数据本身的可靠性进行保证的技术后者指的是对数据的安全防护。

从 Safety 来看，X 轴与 Y 轴是在同一个平面，都是数据库的安全性保障（即高可靠性），分别是 Region 内和 Region 间的高可靠；而 Security 是 Z 轴维度，指的是数据安全，通过例如传输加密、存储加密等能力对数据进行保护。

最后还有斜线是供应链的安全，也就是安可评测的要求。

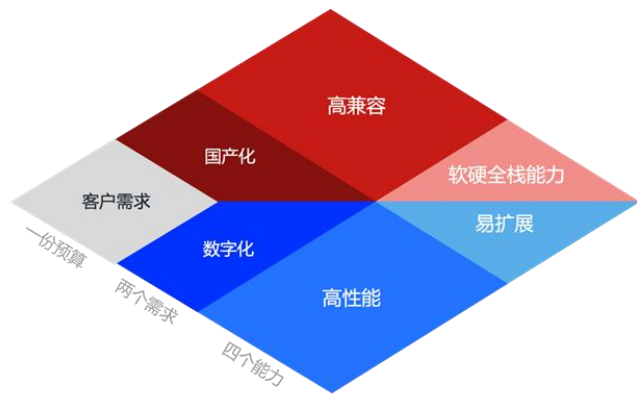
四、数字原生时代的数据库能力需求

对于不同行业，在当前的国产化和数字化“双轮驱动”的情况下，数据库具体要求到什么程度呢，本文进行了统计分析：

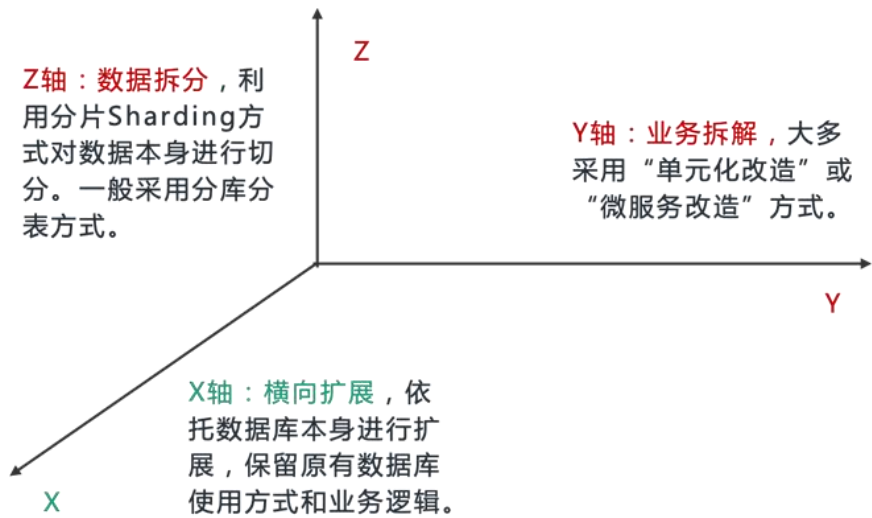


当前国内正处于数字化和深化国产的阶段，对于数字原生组织来说，需要的是高负载、高并发高可靠的数据库，分布式架构是一个更好的选择。对于数据移民来说，目前正在由集中式架构逐步往分布式架构进行转换，这个过程在业界也被称之为“数字化转型”或者“产业数字化”。最后一部分-数据新民行业机构，还在信息化建设和国产化对标替换的状态下，未来才会逐步过渡到高负载的分布式场景中。

平衡新时代的数字化和国产化双重诉求，对数据库的能力下图可以作为借鉴参考：



但是目前国内的数据库厂商，在同时满足这两个诉求上，还有一些工作要做。传统集中式厂商主要在上半部分发力，云和互联网厂商主要在下半部分发力。引入 AFK 方法对数字化转型的主流手段进行分析：



国内目前具有三种数字化转型（数据扩展）的方法：

1. 业务拆解：大多采用容器化手段对业务整体按流程进行“单元化改造”和“微服务改造”。选择此方法的客户以头部金融行业为主，具有强大的技术实力与开发团队。
2. 数据拆分：即分库分表方案，将业务和数据的分布规则进行匹配。采用该方案的以互联网企业、头部央企为主，对业务深入且具备对数据库持续维护升级的能力。
3. 横向扩展：即原生分布式产品级扩展，腰部企业和数据移民企业采用较多的方案，技术团队实力不强，信息中心或科技部门更关注本身业务的组织。

19、从数据库到数据计算系统

一、数据系统的新趋势

随着科学技术的发展，特别是信息技术的发展，人类获取数据、存储数据、分析数据的能力有了划时代的进步。近年来，移动互联网、物联网、5G 等技术持续发展，全球数据圈（Global Datasphere）呈指数级递增，IDC 预测全球数据圈将于 2025 年增长值 175ZB，而中国的数据圈有望于 2025 年爆炸式增长为世界第一。众多围绕数据获取、数据存储、数据分析的应用也如同雨后春笋般地涌现出来。这些事实说明，数据很有可能成为人类科技和经济进步的下一个引爆点。因此人们提出了“数据要素”的

概念，把数据列为和“人力”、“资源”等并列的生产要素，由此可见人们已经对数据的定位产生了本质的变化。

2023 年国家正式成立了国家数据局，负责协调推进数据基础制度建设，统筹数据资源整合共享和开发利用，统筹推进数字中国、数字经济、数字社会规划和建设等，不仅体现了对数据资源的战略性管理和规范化利用的需求，也体现了国家层面对数字经济发展和数据治理的重视。

海量的数据被收集、存储，也为人工智能（AI）的技术的发展提供了前所未有的机遇，各种使用深度学习（Deep Learning）技术和大语言模型（Large Language Model, LLM）的应用也因为海量数据的支持变得更加准确和智能。

此外，云平台的快速发展也为数据系统带来了新的发展机遇。云平台代表了目前最大的计算能力、存储能力和水平扩展能力。云平台为数据系统提供了近乎无限的存储资源和算力资源，使得类似 ChatGPT 这样的应用成为可能。因此越来越多的企业将应用向云上迁移，而越来越多的数据也流向云上。

云计算技术、人工智能（AI）等多种技术的快速发展，数据系统也展现出云原生和智能化的趋势。数据分析技术也在从传统的 BI 向支撑深度学习、大语言模型等新型应用演进。数据系统的走向云原生化和 AI 化是数据系统发展的新趋势。

二、从数据库到数据计算系统

人们通常说的数据库是指关系型数据库管理系统（Relational Database Management System, RDBMS）。除了关系型数据库之外，还有针对各种非结构化数据的数据库，例如文档（Document）数据库、图（Graph）数据库、流式（Streaming）数据库等。企业的数据平台一般由多种数据库联合组成，不同种类的数据库分别处理不同的数据。为了统一管理，人们提出了“企业数据湖”的概念，一般来说，多种数据源和针对不同数据源的处理工具组成了企业数据湖。在数据湖中，数据根据使用的频率分为“冷-温-热”数据，“热”数据一般存储在数据仓库中，在处理“温”数据或者“冷”数据时，一般需要有一个“由湖入仓”的“数据抽取-数据转换-数据加载”（ETL）操作，把数据加载到数据仓库中。

为了打破各种数据种类的边界，降低移动数据的 ETL 操作的代价，人们提出了“湖仓一体”架构，来解决例如事务支持、数据模型化、数据治理等问题。云原生技术和人工智能技术的发展，对数据库、数据湖的架构带来了潜移默化的影响。一种新型的数据架构“数据计算系统（data computing system）”被业界提出。数据和计算是数据计算系统的两个独立子系统，其中数据是核心，计算是产生数据价值的手段。

1.数据子系统

- (1) 元数据（metadata）从用户数据中分离，各个计算任务提供数据的发现、授权、使用等服务。
- (2) 用户数据打破数据类型和数据模式的边界，使用统一的格式存储。



把元数据当作保险箱数字Key数据，把用户数据当做保险箱里的数据黄金。我们只需要对数字Key进行交换，就可以访问保险箱里面的数据。把云上无感知计算当作一堆计算器，需要的时候，根据授权的数字Key，拉到对应的保险箱数据进行计算。

2.计算子系统，针对不同种类的数据使用不同的计算引擎

- (1) SQL 引擎用来执行关系型数据的计算任务
- (2) 向量引擎用来执行与向量数据相关的计算任务
- (3) 图数据引擎用来执行与图数据相关的计算任务
- (4) Python/R 等引擎用来执行 BI/AI 相关的 Python/R 的计算任务

除此之外，数据计算系统还有如下特点：

3.事务支持：对事务的 ACID 支持，可确保数据并发访问的一致性、正确性。

4.数据治理：保证数据完整性，并且具有健全的治理和审计机制。

5.BI 支持：支持直接在源数据上使用 BI 工具，这样可以加快分析效率，降低数据延时。

6.开放性：采用开放、标准化的存储格式提供丰富的 API 支持，因此，各种工具和引擎（包括机器学习和 Python / R 库）可以高效地对数据进行直接访问。

7.支持多种数据类型

8.支持各种工作负载：支持包括数据科学、机器学习、SQL 查询、分析等多种负载类型。

20、自然语言交互

一、背景

数据库系统的初衷是为信息社会提供数据服务，但是使用这些系统需要掌握专门的语言，技术门槛较高。

如图 1 所示，数据库自然语言交互是指用户可以轻松地通过自然语言与数据库进行交互，提交查询和完成操作数据请求，从而简化了数据库的使用，让更多的人能够从中获取有价值的信息。

近年来，随着深度学习和自然语言处理技术的发展，人们对用自然语言与数据库进行交互的需求不断增加。许多研究人员利用机器翻译技术训练通用且具有泛化能力的翻译模型，以实现这一过程。

最近，大语言模型（LLM）技术的跃进为自然语言与数据库的交互形式开辟了新的道路。与传统的有监督学习训练获得的翻译模型不同，这类模型展现出了令人瞩目的涌现能力，使模型能够无缝地将自然语言翻译为所需的交互输出，而无需专门的模型训练过程。这种范式转变正在革新数据库自然语言交互的方式，推动数据库系统的普及和应用，并促进人们更高效地利用数据资源获取有益信息。

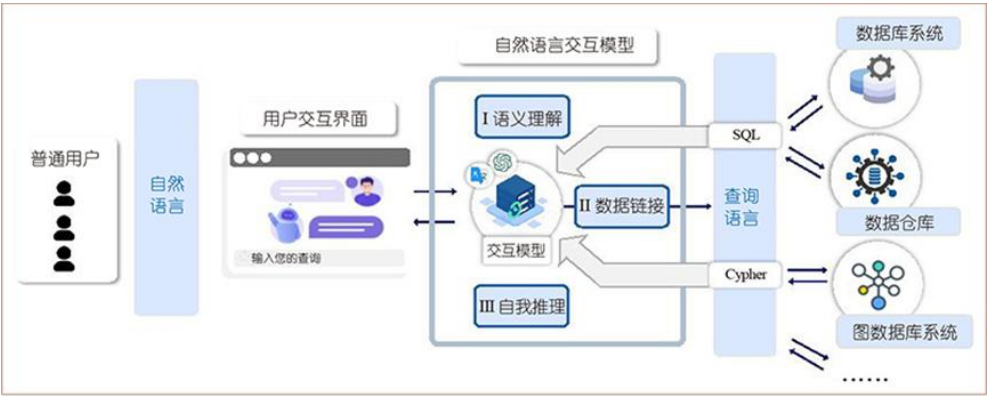


图 1：数据库自然语言交互方式

二、数据库自然语言交互技术

目前，数据库自然语言交互技术主要采用机器学习方法和大语言模型方法，这些方法具有智能和泛化能力，可以为用户提供智能的数据库访问方式。

（一）基于语言翻译的自然语言交互

在进行自然语言交互的过程中，语言翻译模型利用序列到序列（Seq2Seq）网络框架，将目标查询语言（如 SQL）视为一种特殊的“小语种”，并通过有监督训练在标记数据上学习翻译能力，完成从自然语言文本到目标查询语言的转化。

正如图 2 中所示，语言翻译模型首先使用 Seq2Seq 框架中的编码器网络模型对自然语言查询进行编码，以表征自然语言查询的语义信息。同时，为了更好地理解自然语言，模型也需要对数据视图的语义信息进行有效编码，从而获得自然语言和数据视图的统一编码表示（可以利用编码器本身，或借助额外的网络模型）。然后，语言翻译模型将编码过程中获得的特征表征作为输入，借助 Seq2Seq 框架中的解码器网络模型，按序（即符合人类读写顺序）生成目标查询。

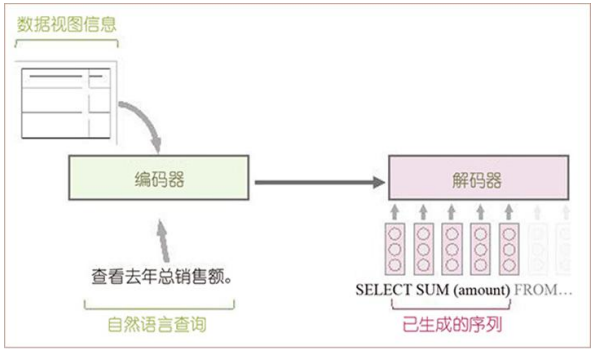


图 2：使用语言翻译模型进行自然语言交互的示例

（二）基于大语言模型的自然语言交互

大型语言模型的应用作为数据库自然语言接口引起了广泛关注。这些模型通过学习大量数据，展现出了卓越的复杂推理能力，并能够在无需进行特定数据微调的情况下展现出色的性能。举例来说，在公开基准测试数据集 Spider 中，利用大型语言模型实现的翻译准确率已经从过去的 53.5% 提高至 85.3%，出现了质的飞跃。

由于其强大的语言理解能力和文本生成能力，大型语言模型能够直接作为数据库的自然语言交互模型使用。在交互过程中，用户可以通过零样本提示和少样本提示方式使用大型语言模型。零样本提示是指使用简单的任务描述来让模型实现翻译过程；而少样本提示则是通过提供少量翻译示例或者推理示例来让大型语言模型“模仿”示例的输入输出信息，从而实现翻译过程。

在 2023 年 12 月《中国计算机学会通讯》月刊上，复旦大学的学者对当前可用的大型语言模型作为自然语言交互模型的性能进行了初步分析。分析中对主流的基于 OpenAI GPT 的一系列模型（包括 CodeX Davinci、GPT-3.5-turbo 和 GPT-4 模型）、基于 GPT-4 改进的 DIN-SQL 模型、基于 Google Pathways 语言模型（PaLM）的 SQL-PaLM 模型，以及两个传统翻译模型 LGESQL 和 RESDSQL 的性能进行了对比。该对比采用跨领域评测数据集 Spider，包含 10181 条自然语言查询和 5693 条不同的

复杂 SQL 查询，涵盖 206 个数据库，覆盖了 138 个不同领域。对比实验中采用 Spider 中定义的“查询执行准确率”和“完全匹配准确率”两个评测指标对模型进行评估。其中，查询执行准确率指标表示翻译得到的查询在数据库执行后输出结果的准确性，完全匹配准确率是指翻译得到的查询在语法上完全一致的准确性。

表 1 呈现了各种模型的性能对比情况。由于 Spider 测试集未公开，需要将其提交到评估服务器进行评测，因此只有少数模型获得了评测结果。从对比结果可以得出以下结论：

- 1. 总体而言，大型语言模型在查询执行准确率方面表现优于传统翻译模型；而在完全匹配准确率方面，传统翻译模型表现更出色。
- 2. 大型语言模型通过零样本提示方式成为数据库自然语言接口，目前还未实现很好的翻译性能，其性能效果总体比少样本提示方式差。例如，对于 ChatGPT（即 GPT-3.5-turbo）模型，在 Spider 验证集数据上通过零样本提示方式仅能实现 38.4% 的完全匹配准确率，比少样本提示方式的结果低 13.1%。
- 3. 通过少样本提示方式，大型语言模型能够根据提供的少量样本进行更好地理解 and 推理，实现与传统翻译模型相近的性能效果。此外，通过设计更好地提示内容，大型语言模型能够展现超越现有传统翻译模型的性能效果。例如，DINSQL 模型通过设计一系列精细化的提示信息，在 Spider 测试集数据上实现了 85.3% 的查询执行准确率，比目前最好的 RESDSQL 传统翻译模型性能高 5.4%。

	模型	Spider验证集		Spider测试集	
		查询执行 准确率	完全匹配 准确率	查询执行 准确率	完全匹配 准确率
零样本提示	CodeX Davinci	47.5%	—	—	—
	GPT-3.5-turbo	60.1%	38.4%	—	—
	GPT-4	64.9%	40.4%	—	—
	SQL-PaLM	76.0%	—	—	—
	DINSQL	64.9%	40.4%	—	—
少样本提示	CodeX Davinci	61.5%	50.2%	—	—
	GPT-3.5-turbo	63.5%	51.5%	—	—
	GPT-4	67.4%	54.3%	—	—
	SQL-PaLM	77.3%	—	—	—
	DINSQL	74.2%	60.1%	85.3%	60%
有监督学习	LGESQL	36.3%	75.1%	34.2%	72.0%
	RESDSQL+T5-3b	84.1%	80.5%	79.9%	72.0%

表 1 自然语言交互模型性能对比结果（摘自 2023 年 12 月《中国计算机学会通讯》）

三、自然语言交互中的挑战问题

应用大语言模型于数据库自然语言接口仍然面临着一系列产业化落地的挑战。

首先，翻译性能的提升一直是长期存在的问题。由于自然语言与底层数据在语义表达上存在鸿沟，大语言模型仅凭借世界知识无法解决领域知识的语义理解问题。举例来说，SQL 中的连接（JOIN）操作产生的新数据表可能存在特殊的语义信息，仅利用世界知识无法准确地理解。此外，大语言模型在自然语言到数据库查询的翻译过程中，容易产生错误，特别是涉及数据值预测时。虽然大语言模型学习获得的世界知识有助于理解自然语言中的语义，但在特定数据库的数据结构下，大语言模型在自然语言到数据实体映射的过程中容易出现错误。此外，在某些情况下，比如处理复杂的自然语言查询时，大语言模型会生成一些实际并不存在的数据实体，产生大模型幻觉问题。

另外，模型的复杂性和计算资源需求也是一个重要挑战。大语言模型通常需要大量的计算资源和存储空间，限制了它在实际生产环境中的可行性。另外，直接使用数据库数据进行模型参数微调虽然可以实现特定数据库的适配、进一步提高大语言模型的交互性能，但不适合数据快速动态变化的场景，且微调的代价较高，需要消耗大量计算资源和时间。

四、消除大模型幻觉的有效途径——RAG

尽管大语言模型的能力令人印象深刻，但它们并非毫无瑕疵。这些模型可能会产生误导性的“幻觉”，依赖的信息可能过时，处理特定知识时效率不高，缺乏专业领域的深度洞察，同时在推理能力上也有所欠缺。在数据库自然语言查询到 SQL 的翻译过程中，大模型幻觉问题的存在会使翻译生成的 SQL 语句中出现数据库中不存在的实体，导致翻译错误。

在实际应用中，数据需要不断更新以反映最新的发展，生成的内容必须是透明可追溯的，以便控制成本并保护数据隐私。因此，简单依赖于这些“黑盒”模型是不够的，我们需要更精细的解决方案来满足这些复杂的需求。正是在这样的背景下，检索增强生成技术（Retrieval-Augmented Generation, RAG）应运而生。RAG 通过在大语言模型生成答案之前，先从数据库中检索相关信息，然后利用这些信息来引导生成过程，能极大提升生成内容的准确性和相关性。RAG 有效地缓解了幻觉问题，提高了知识更新的速度，并增强了内容生成的可追溯性，使得大语言模型在实际应用中变得更加实用和可信。

一个典型的 RAG 案例如图 3 所示。如果我们向 ChatGPT 询问 OpenAI CEO Sam Altman 在短短几天内突然解雇随后又被复职的事情。由于受到预训练数据的限制，缺乏对最近事件的知识，ChatGPT 则表示无法回答。RAG 则通过从外部知识库

检索最新的文档摘录来解决这一差距，它获取了一系列与询问相关的新闻文章，连同最初的问题一起合并成一个丰富的提示，使 ChatGPT 能够综合出一个有根据的回答。

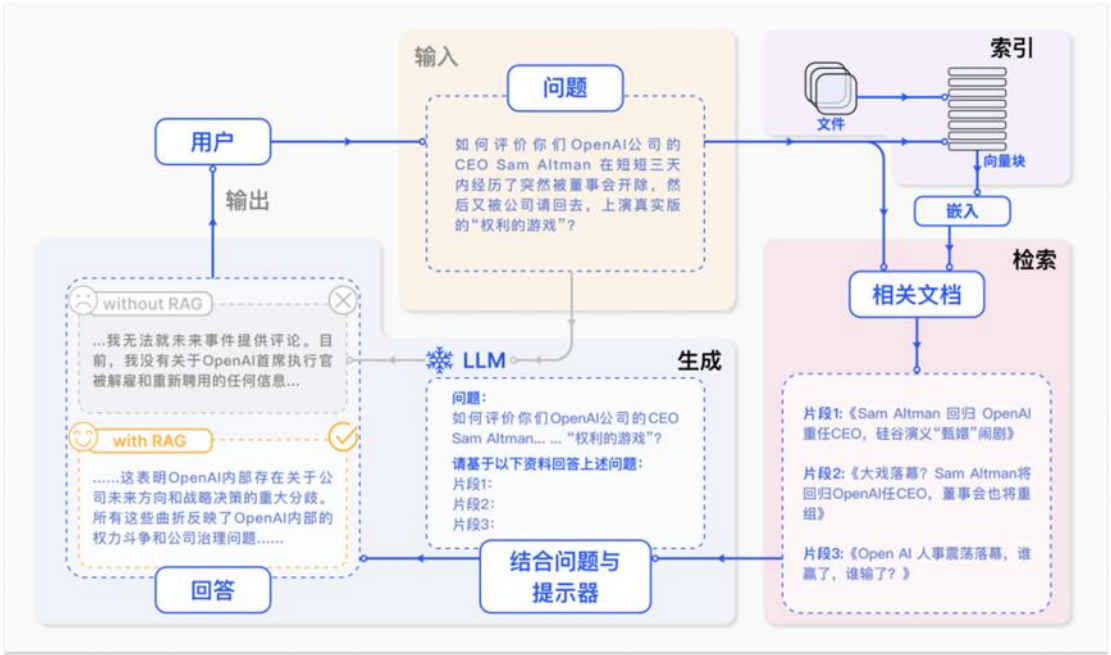


图 3： RAG 技术在自然语言交互中的案例

五、RAG 技术的发展

RAG 概念于 2020 年首次提出，并随后迅速发展。在早期的预训练阶段，研究的重点是如何通过预训练模型注入额外的知识，以增强语言模型的能力。随着 ChatGPT 的推出，对于使用大语言模型进行深层上下文学习的兴趣激增，推动了 RAG 技术在研究领域的快速发展。随着大模型潜力的进一步开发，旨在提升模型的可控性并满足不断演变的需求，RAG 的研究逐渐聚焦于增强推理能力，并且也探索了在微调过程中的各种改进方法。特别是随着 GPT-4 的发布，RAG 技术经历了一次深刻的变革。研究重点开始转移到一种新的融合 RAG 和微调策略的方法，并持续关注对预训练方法的优化。

在 RAG 的技术发展过程中，从技术范式角度可以将其总结成如图 4 所示的几个阶段：

（一）朴素 RAG 阶段

在图 3 案例中展示了经典的 RAG 流程，也被称为朴素 RAG。主要包括三个基本步骤：

1. 索引 — 将文档库分割成较短的片段，并通过编码器构建向量索引。
2. 检索 — 根据问题和片段的相似度检索相关文档片段。

3. 生成 — 以检索到的上下文为条件，生成问题的回答。

(二) 进阶 RAG 阶段

朴素 RAG 在检索质量、响应生成质量以及增强过程中存在多个挑战。进阶 RAG 范式随后被提出，并在数据索引、检索前和检索后都进行了额外处理。通过更精细的数据清洗、设计文档结构和添加元数据等方法提升文本的一致性、准确性和检索效率。在检索前阶段则可以使用问题的重写、路由和扩充等方式对齐问题和文档块之间的语义差异。在检索后阶段则可以通过将检索出来的文档库进行重排序避免 “Lost in the Middle” 现象的发生，或者通过上下文筛选与压缩的方式缩短窗口长度。

(三) 模块化 RAG 阶段

随着 RAG 技术的进一步发展和演变，新的技术突破了传统的朴素 RAG 检索—生成框架，最终引入了模块化 RAG 的概念。在结构上它更加自由的和灵活，引入了更多的具体功能模块，例如查询搜索引擎、融合多个回答。技术上将检索与微调、强化学习等技术融合。流程上也对 RAG 模块之间进行设计和编排，出现了多种的 RAG 模式。模块化 RAG 并不是突然出现的，三个范式之间是继承与发展的关系。进阶 RAG 是模块化 RAG 的一种特例形式，而朴素 RAG 则是进阶 RAG 的一种特例。

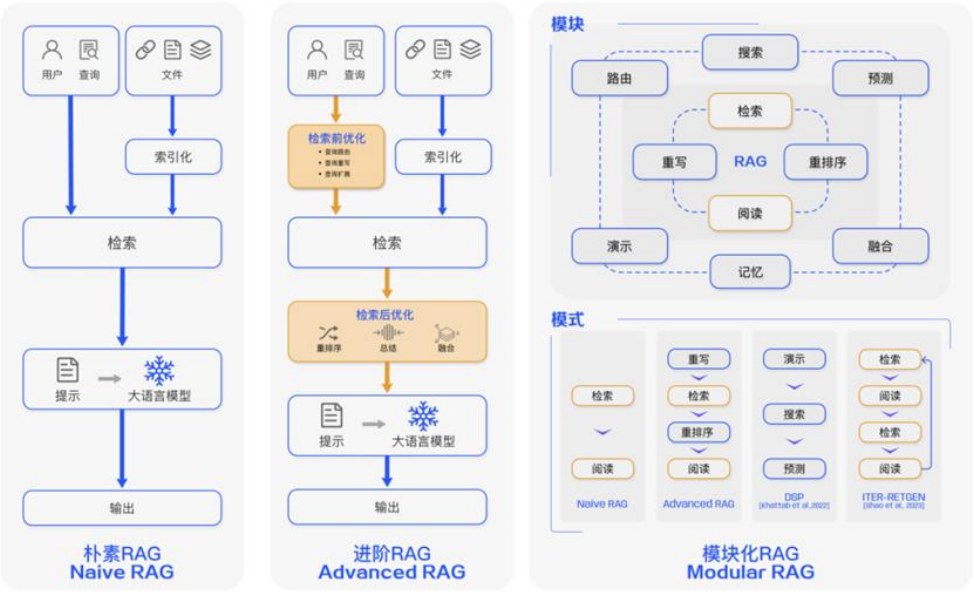


图 4：RAG 发展阶段的范式对比

六、RAG 与微调技术的融合

除了 RAG，微调（Fine-tuning，FT）也是大语言模型的主要优化手段。微调主要用于改进预训练模型的性能，可以分为微调所有层、微调顶层、冻结底层、逐层微调和迁移学习等多种微调方式。目前前沿的微调方法主要包括以下几种：

1.Prompt tuning

Prompt tuning 是一种参数高效性微调方法，它固定模型的前馈层参数，仅更新部分 embedding 参数来实现低成本的大模型微调。应用这种微调方法的侧重点在于精心制作可以指导预训练模型生成所需输出的输入提示或模板。

2.Prefix Tuning

Prefix Tuning 通过在预训练模型的输入端添加一个特定的前缀来实现微调。这个前缀可以是一个特定的标记序列，也可以是一个特定的模式，它可以指导模型在微调过程中更好地适应特定的任务。通过在输入端添加前缀，Prefix Tuning 可以在不改变预训练模型参数的情况下，针对特定任务进行微调，从而提高模型在该任务上的性能。

3.P-tuning v2

P-tuning v2 是 P-tuning 的改进版本。P-tuning v2 通过对模型的参数进行微调，以提高模型在特定任务上的性能。与传统的微调方法相比，P-tuning v2 使用了更加精细的参数调整策略，可以更好地适应不同的任务和数据集。此外，P-tuning v2 还引入了一种新的参数初始化方法，可以更好地避免模型陷入局部最优解。

4.Lora

Lora（Language Representation with Attention）是一种基于注意力机制的大语言模型微调方法。它通过在预训练模型上添加一个额外的任务特定的输出层，并使用任务特定的数据集进行微调，从而使模型能够适应特定的任务。这种方法通常能够在少量的标注数据上取得很好的效果，因此在实际应用中得到了广泛地应用。

5.RLHF——人类反馈强化学习

RLHF（Reinforcement Learning with Human Feedback）是一种结合强化学习和人类反馈的大模型微调方法。在 RLHF 中，模型首先通过强化学习算法进行预训练，然后将其应用于实际环境中。在实际环境中，人类可以提供反馈，指导模型进行微调和改进。这种方法可以帮助模型更快地适应实际环境，并且可以利用人类的专业知识和经验来指导模型的学习过程。通过结合强化学习和人类反馈，RLHF 可以提高模型的性能和适应能力。

6.DeepSpeed ZeRO——零冗余优化

DeepSpeed ZeRO 是一种用于大规模模型微调的方法，它通过使用零内存优化技术（Zero Redundancy Optimizer）来减少内存占用和加速训练过程。这个方法把模型参数分成多个部分，然后在每个部分上执行独立的优化和梯度累积，从而减少了内存占用。这样就可以在内存有限的情况下训练更大的模型，同时也加速了训练过程。DeepSpeed ZeRO 还提供了一些其他优化技术，比如混合精度训练和动态图优化，以进一步提高训练效率。

特点对比	检索增强生成 (RAG)	微调 (FT)
知识更新	RAG直接更新检索知识库，保持信息最新，模型无需频繁的重新训练，适合动态数据环境。	FT存储静态数据，知识与数据更新需要重新训练。
外部知识	RAG擅长利用外部资源，非常适合文档或其他结构化/非结构化数据库。	虽然FT可以对大语言模型进行微调以对齐预训练学到的外部知识，但对于频繁更改的数据源来说可能不太实用。
数据处理	RAG对数据加工和处理的要求低。	FT依赖高质量数据集，有限的数据集可能不会产生显著性能提升。
模型风格	RAG主要关注信息检索，擅长整合外部知识，但可能无法完全定制模型的行为或写作风格。	FT允许根据特定的语气或术语调整大语言模型的行为、写作风格或特定领域的知识。
可解释性	RAG通常可以追溯到特定数据源的答案，从而提供更高等级的可解释性和可溯源性。	FT就像黑匣子，并不总是清楚模型为何会做出这样的反应，具有相对较低的可解释性。
计算资源	RAG需要高效的检索策略和大型数据库相关技术。另外还需要保持外部数据源集成以及数据更新。	FT需要准备和整理高质量的训练数据集、定义微调目标以及相应的计算资源。
延迟和实时要求	RAG需要进行数据检索，可能会有更高延迟。	经过FT的大语言模型无需检索即可响应，延迟较低。
减少幻觉	RAG本质上不太容易产生幻觉，因为每个回答都建立在检索到的证据上。	FT可以通过将模型基于特定领域的训练数据来减少幻觉。但当面对不熟悉的输入时，它仍然可能产生幻觉。
道德和隐私问题	RAG的道德和隐私问题来源于从外部数据库检索的文本。	FT的道德和隐私问题则因为模型的训练数据存在敏感内容。

表 2：检索增强生成技术与微调技术的对比

RAG 与 FT 具有各自的特点，如表 2 所示。根据对外部知识的依赖性和模型调整要求上的不同，各自有适合的场景。RAG 就像是为模型提供一本定制的信息检索教科书，非常适合特定的查询。FT 则类似于学生随着时间内化知识，更适合模仿特定的结构、风格或格式。FT 可以通过增强基础模型知识、调整输出和教授复杂指令来提高模型的性能和效率。然而，它不太擅长整合新知识或快速迭代新的用例。RAG 和 FT 并不是相互排斥的，它们可以相互补充，联合使用更有希望产生最佳性能。

七、产业趋势与实践

未来的数据库将迎来自然语言交互的主流趋势。随着人工智能（AI）和大模型技术的迅速发展，数据库自然语言交互将具备更强大的语言理解能力，从而实现更加自然和无缝的对话交流。我们预测未来数据库自然语言交互将呈现以下两大趋势。

首先，结合检索增强生成技术、微调技术和提示工程等多种优化手段，来提升大语言模型在数据库自然语言交互中的翻译效果。就目前的技术水平来看，把自然语言翻译成 SQL 还有很大的改进空间，而要实现这种改进不太可能只靠一种大模型的优化手段。

其次，利用数据库计算引擎与查询引擎的特长来减轻大语言模型的负担。比如说，分布式数据库可以利用自己的分布式计算和存储系统来加速大型模型的微调和推理。它还可以在多个节点上平衡大型模型的计算资源需求，实现任务、数据和模型的并行处理。任务并行就是把一个大任务分成小任务，然后同时执行，这样可以加快任务的执行速度。数据并行是把数据分成多部分，然后同时处理，最后合并结果得到最终结果。模型并行是把一个大模型分成小模型，然后同时训练或推理，最后合并参数得到最终模型。

贝格迈思是国际领先的数据科技驱动行业应用创新者，致力于利用新型硬件革新和软件技术进步，融合国际领先的机器学习算法和大模型技术框架，打造真正实现算力革新的智能数据库 AiSQL。AiSQL 是一个能够统一管理数据和 AI 模型的工具，可以在多种深度学习框架下支持大语言模型的前沿微调技术和模块化检索增强生成技术。作为分布式数据库，AiSQL 利用自身的分布式存储优势，为模型和应用数据提供高可用的分布式多副本存储机制。此外，AiSQL 能够在进程内与 AI 算法模型交换数据，并利用分布式部署支持大语言模型的微调和推理。通过前沿的并行处理方式，AiSQL 保证了大模型应用的高效性和可扩展性。AiSQL 为企业提供了一种高效、可靠的数据和 AI 模型管理解决方案，能够帮助企业提升业务效率 and 创新能力。

21、自动驾驶数据库

近年来，随着云计算技术的迅速发展，在云环境上构建、部署并提供数据库访问能力的云数据库系统快速崛起，在商业应用上正在改变传统数据库应用的格局。无论是国内还是国际，云数据库成为未来数据库系统发展的主流已成为必然趋势。为适应云环境下部署、应用数据库的需求，云原生数据库（Cloud Native Database）系统应运而生，由于云原生数据库系统必须从系统内部进行存储与计算的解耦，通过存算分离的架构满足系统对于高弹性扩展的需求以充分融合和释放云技术优势，其在设计时会比本地部署（On-Premise）的传统集中式数据库面临更大的权衡维度和优化空间。云原生数据库系统不但要适配不同的云服务商，也需要在高性能和代价的权衡下选择合适的硬件资源，同时还需要在多种存储引擎中选择合适的存储形态，最终结合云服务 SLA 的保证，满足应用的性能约束。由于云环境软硬件的复杂性以及云原生数据库的环境敏感性，采用人工方式进行系统优化很难产生正确的决策，即使是次优的决策也会因为云环境的复杂性而放大性能损失和开销，而且这种手动决策如果产生错误的选择，相比于传统数据库往往会产生更加严重甚至是灾难性的负面影响。此外，对于云环境的数据库管理员而言，需要面对成千上万个并发运行在云环境下的数据库实例，人工为之配置合适的资源

几乎成为不可能完成的任务。对云负载的研究表明：很多云数据库应用没有很好地利用系统的配置设置、模式设计或者存取结构。解决云原生数据库在资源利用、性能优化以及提供良好的用户服务方面的有效途径是研究一种具有能够自动选择改善性能的动作的能力，自动选择何时应用这些动作的能力以及自动从其动作与决策中学习的能力的新型数据库管理系统。这里动作的含义包括：对数据库物理设计的改变、对数据库旋钮配置的改变、对数据库物理资源的改变等，该类系统即为自动驾驶云原生数据库系统（Self-driving Cloud Native Database）。

自动驾驶数据库系统是指能够针对不同的硬件、软件、应用场景、负载特征、性能要求以及设计约束等外部多重关联特征，根据现状以及对未来的预测，不再依赖人工介入实现完全自动化的系统优化，在性能、易用和稳定性上实现动态平衡，进而达到智能的自治。事实上，数据库系统的优化一直以来都是研究的核心与焦点，不管是从外部多重关联维度的离散或者连续参数空间中，还是从内部可生成或选择操作的状态空间中，为数据库系统找到一个匹配外部多重关联特征的性能最优解都是极其复杂的 NP 问题。数据库系统从诞生之初，就试图通过经验法则、专家辅助以及被动式的自适应技术解决数据库系统的性能优化问题，但是在海量数据的复杂场景下实际应用效果并不理想。近年来，在人工智能技术高速发展的驱动下，结合人工智能对数据库进行调优的 AI4DB 的研究引发了研究者大量的关注，但这些研究尚处于初期阶段，且大多数研究内容仅停留在原有的数据库架构上利用机器学习技术对部分组件进行有限的局部优化。因此，研究成果的泛化能力、表现能力、适用能力以及优化效果都难以满足实际应用需求。实现真正自动驾驶系统，需要从全新的系统架构出发，对计算引擎、存储引擎以及所有可规则化组件进行智能化的重构，是一项极具挑战性的工作。目前，仅有少数研究如 NoisePage 提出了类似的思想。

减轻或摆脱数据管理的人力负担一直是数据库管理技术发展的主要推动力之一。关系数据模型以及描述性查询语言就是典型的例证，在关系系统中，开发者只需要写出一个需要存取什么数据的查询，数据库管理系统负责提供高效地存储与检索数据的方式，进一步地，希望数据库能够提供管理与优化数据库的整体策略，这就是 70 年代的自适应数据库系统概念，自适应数据库管理系统主要聚焦于物理数据库设计以及索引的选择，其主要工作方式是：系统收集应用存取数据的模式信息，然后基于代价模型搜索改善性能的策略。90 年代后开始了自动调优数据库系统的研究并一直持续到现在，最初主要集中于研究辅助工具帮助 DBA 选择最优索引、物化视图以及分区模式等，2000 年代开始了自动旋钮设置的研究，这些旋钮允许 DBA 控制数据库运行时行为的各个方面，这些工具的使用都是手动过程：DBA 提供样本负载，调优工具根据这些负载提供一个或多个调优建议，DBA 决定是否采纳这些建议。自适应调优是数据库系统实现自我管理的一种透明化技术途径，但是经验法则驱动和优化方式以及感知式的优化策略往

往难以应对复杂环境下的动态调整需求，特别是当云原生数据库需要面对复杂的云环境和高并发的异构负载保持符合应用需求的性能时，其面临的是比以往数据库更多层的优化空间和更高维的调整旋钮，探索和研究自动驾驶优化技术就显得尤为重要和紧迫。近年来随着大数据驱动的人工智能特别是机器学习技术的迅猛发展，由于其较强的应用场景感知和自我学习能力，在数据库性能优化研究领域呈现出了巨大的潜力。

近年来各种数据库原型系统和产品在自动优化方面都开展了不少的研究，但现有技术存在很多局限性，难以在实际应用场景中有效落地，究其原因主要是云原生数据库具有以下几个独特性质：

1.基于学习的优化技术泛化能力欠缺：云原生数据库系统需要面对复杂的云环境（包括不同的云服务商、硬件规格、服务水平、资费标准等）和异构多变的工作负载（包括 TP、AP 或者混合以及不同的读写特征等），目前的优化技术在实际应用中适用性有限，效果不佳，重复训练时间长，因此需要研究具有高泛化能力的自动驾驶优化技术。

2.基于重构的系统性自动驾驶优化难度极大：当前的优化技术侧重以最小的实现工作量非侵入的添加优化技术，或者以简单的接口添加即插即用的优化组件。但是，云原生数据库面临的优化空间巨大、系统复杂度极高，因此只有从系统内部进行模块化拆分并通过细粒度的优化建模和适配，才有机会获得预期的优化效果，但是实现复杂度极高。

3.人工智能技术与数据库技术的匹配度不足：近年来随着人工智能特别是深度学习、强化学习以及深度强化学习的快速发展，为实现自动驾驶的数据库优化提供了理论上的基础。但是即使是在重构系统中针对单元化的模块进行优化，也面临寻找最优配置、最优计划等 NP-hard 难题。所以，需要实现人工智能技术与数据库技术的高度融合才能产生实际价值。

机器学习是目前自动驾驶云原生数据库研究的主要手段，但并不是唯一的技术途径，感知式或者启发式的方法在其中都仍然具有重要的地位，这些实现自动驾驶的技术可以相互融合地使用。其次，虽然目前的研究大多聚焦在查询优化、配置优化等几个方面，但是理论上，数据库中任何具有经验性规则的模块本质上都可以利用自动驾驶技术完成自适应的优化。自动驾驶云原生数据库的研究不是简单地将人工智能或感知技术应用于系统的某个组件上就能带来预想的效果，特别是叠加云计算和 HTAP 异构负载后，很难通过简单的修补式优化到达高等级（L5 级）的自动驾驶。因此全面解耦数据库自动驾驶的组件，研究从生成、选择、执行、反馈、调整等一系列贯穿整个数据库生命期的自动驾驶技术就成为一种必然的选择，这种全局式优化技术的研究刚刚起步。

云计算技术的广泛使用，使国产数据库系统的发展迎来了新的发展机遇。以提供数据库云服务为目标的云原生数据库管理系统的研发在国内取得了显著的进展，并开始在自动驾驶研究方向开展了前沿布局，为数据库的国产替代奠定了良好的基础，如 GaussDB、PolarDB、YaoBase 等，已经具备了初级的自动驾驶潜力。但是，随着数据库业务向云端迁移，对复杂应用环境以及运行环境提供良好的动态适应性是云原生数据库管理系统面临的重大挑战。局部、单一、被动的传统自适应调优技术已经无法应对复杂多变的应用需求及系统环境，发展具有自动管理、自主优化、自我修复、自主学习等特征的完全自动驾驶数据管理技术才能满足对于云原生数据库系统在稳定性、适用性和易用性上的刚性需求。

第三章：数据库新兴硬件及生态产品

1、处理器

一、概述

构建高性能、高可靠、高并发的数据库已成为企业推动数字化转型的重要策略。随着人工智能（AI）、深度学习（DL）、即时分析等应用的不断创新，需要处理与分析的数据出现爆发式增长。企业对于数据库性能的要求也与日俱增，以更充分地挖掘数据价值。要提升数据库性能，不仅依赖于软件层面的创新与优化，还需要尽可能消除 CPU 层面的瓶颈。

第四代/第五代英特尔®至强®可扩展处理器是构建高性能数据库的理想选择，其不仅具备强大的核心性能，而且配备更大三级缓存、更快内存和英特尔®数据分析引擎，能够充分释放数据库的性能潜力，加速企业数据战略的实施。英特尔与南大通用、星环科技、PingCAP、腾讯云等合作伙伴携手，构建了基于第四代/第五代英特尔®至强®可扩展处理器的数据库解决方案，并证明其性能足以满足企业数字化转型所需。

二、基于处理器的数据库解决方案

对于数据库应用而言，第四代/第五代英特尔®至强®可扩展处理器提供了强大的性能基础，这不仅源于处理器的硬件性能，也源于处理器内置的英特尔®存内分析加速器（英特尔® IAA）、英特尔®高级矩阵扩展（英特尔® AMX）、英特尔®高级矢量扩展 512（英特尔®AVX-512）等硬件特性。这些特性通过搭配英特尔®OneAPI DPC++/C++编译器等软件工具，有助于提升数据库性能，帮助用户从数据中更快获取洞察，有效提升生产力。

（一）第四代/第五代英特尔至强可扩展处理器

第四代英特尔®至强®可扩展处理器通过创新架构增加了每个时钟周期的指令，每个插槽多达 60 个核心，支持 8 通道 DDR5 内存，有效提升了内存带宽与速度，并通过 PCIe 5.0（80 个通道）实现了更高的 PCIe 带宽提升。第四代英特尔®至强®可扩展处理器提供了出色性能和安全性，可根据用户的业务需求进行扩展。借助内置的加速器，用户可以在 AI、分析、云和微服务、网络、数据库、存储等类型的工作负载中获得优化的性能。

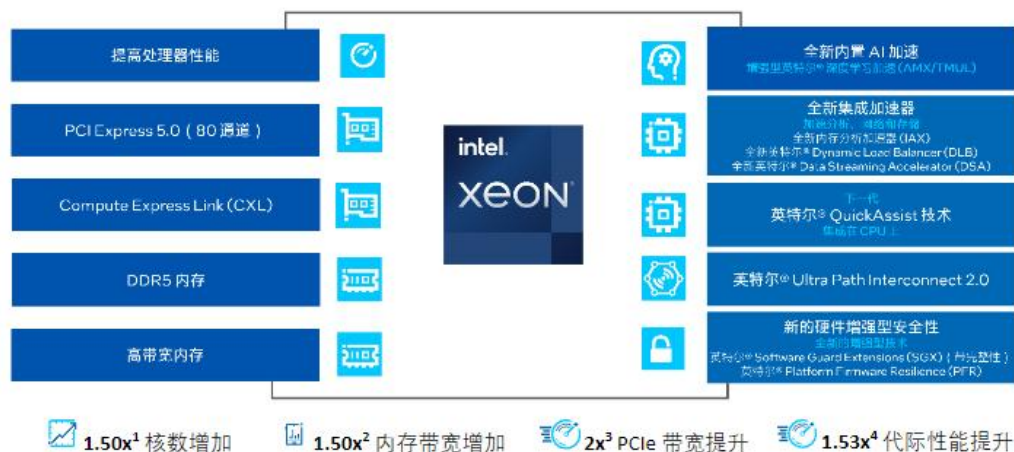


图 1.: 第四代英特尔®至强®可扩展处理器为数据中心提供多种优势

第五代英特尔®至强®可扩展处理器拥有更可靠的性能，更出色的能效。它在运行各种工作负载时均可实现显著的每瓦性能增益，在 AI、数据中心、网络和科学计算的性能和总体拥有成本(TCO)方面亦有更出色的表现。相较上一代产品，第五代英特尔®至强®可扩展处理器可在相同功耗范围内提供更高的算力和更快的内存。此外，它与上一代产品的软件 and 平台兼容，因此部署新系统时可大大减少测试和验证工作。



图 2: 第五代英特尔®至强® 扩展处理器具备更强大性能

第四代/第五代英特尔®至强®可扩展处理器内置了英特尔®IAA、英特尔®AMX 等加速器，以及英特尔®AVX-512 指令集，使其在数据库应用中有更佳的性能表现。

1.英特尔® In-Memory Analytics Accelerator(英特尔® IAA)

在数据库应用中，常需要对海量数据进行压缩存储，从而减少存储数据所需的空间，降低相应的成本投入。这一过程会带来较高的性能开销，英特尔®IAA 能够提升数据压缩率，同时降低性能开销。英特尔®IAA 是一款硬件加速器，结合分析原始函数，能够提供出色的吞吐量压缩和解压缩性能。它主要针对大数据和内存分析数据库等应用程

序，以及内存页压缩等应用程序透明用途，能够在分析查询处理期间过滤数据。该加速器支持零压缩等轻量级压缩方案以及霍夫曼编码和 Deflate 等较重的压缩算法。对于 Deflate 格式，它支持对压缩流进行索引，以实现高效的随机访问。第五代英特尔®至强®可扩展处理器同样集成了英特尔®IAA 加速器，能够提供出色的数据吞吐压缩和解压缩性能。

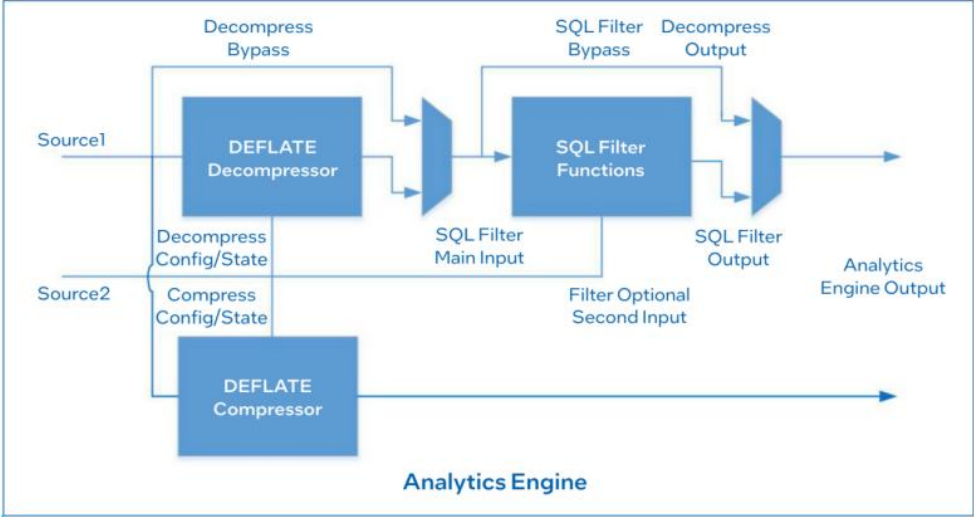


图 3：英特尔® IAA 加速流程

2.英特尔® Advance Matrix Extensions（英特尔® AMX）

在向量数据库等场景，常会涉及向量检索、海量特征匹配等操作，对于性能有着较高的要求。英特尔®AMX 引入了一种用于矩阵处理的新框架（包括了两个新的组件，一个二维寄存器文件，其中包含称为“tile”的寄存器，以及一组能在这些 tile 上操作的加速器），从而能高效地处理各类 AI 任务所需的大量矩阵乘法运算，提升其在训练和推理时的工作效能。例如在向量检索的过程中，如存在 n 个 batch 任务，进行相似度计算时就需要对 n 个输入向量 x 和 n 个数据库中向量 y 进行比对，其中的距离计算会产生大量的矩阵乘法，而英特尔®AMX 能够针对这一场景实现有效加速。

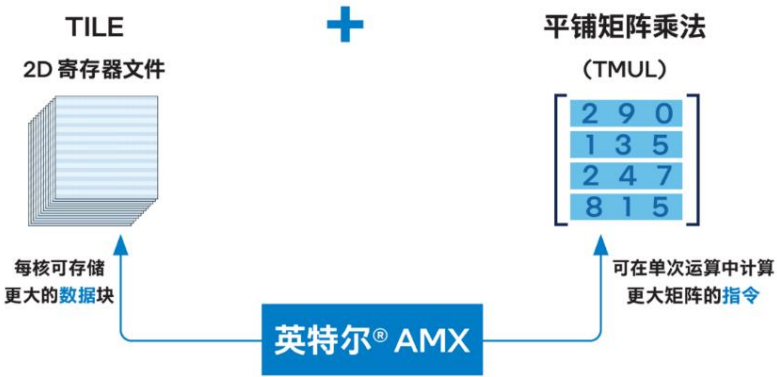


图 4：英特尔®AMX 架构由 2D 寄存器文件 (TILE) 和 TMUL 组成

3. 英特尔® Advanced(英特尔® AVX-512)指令集

在向量数据库等场景，并行计算能力成为数据库性能的重要瓶颈。作为一种单指令多数据（Single Instruction Multiple Data, SIMD）指令集，英特尔® AVX-512 在密集型计算负载中有着得天独厚的优势。得益于其 512 位的寄存器宽度和两个 512 位的融合乘加（Fused Multiply Add, FMA）单元，指令集能并行地执行 32 次双精度、64 次单精度浮点运算，或操作 8 个 64 位和 16 个 32 位整数。

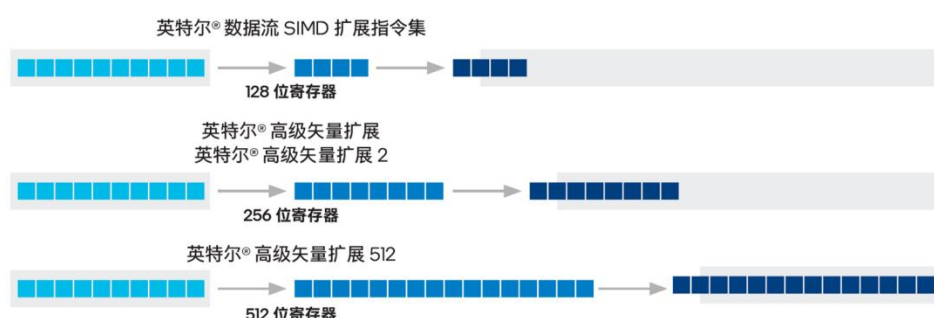


图 5：英特尔® AVX-512 指令集具备 512 位寄存器

（二）软件工具

第四代/第五代英特尔® 至强® 可扩展处理器结合英特尔® OneAPI DPC++/C++编译器等英特尔软件工具，能够进一步提升数据库性能。

1. 英特尔 OneAPI DPC++/C++编译器

英特尔 oneAPI DPC++/C++编译器适用于并行编程（parallel programming）程序，提供跨 CPU 和加速器的生产力和性能。利用该编译器，团队可以结合链接时优化（LTO）和配置文件引导优化（PGO）的方法，帮助用户构建高性能数据库。

2. 英特尔 FMAL 加速库

在海量向量数据处理时，暴力搜索有着非常多的使用，但这一场景对算力需求非常高，因此性能优化极为必要。作为针对向量暴力搜索场景开发的算法库，英特尔® FMAL 在英特尔® AVX-512 和英特尔® AMX 的加持下，能对相似度计算进行加速并提供了相似度计算和 top-K 查询的 API 接口。值得一提的是，英特尔® FMAL 能与英特尔® AMX 结合，对 INT8 数据类型的性能实现进一步优化。同时，英特尔® FMAL 还能在多线程并发下对处理器资源进行合理地调配，以便让用户充分挖掘最新处理器所具

备的多核心优势。除此之外，加速库也提供了对内存的非一致内存访问架构（Non Uniform Memory Access, NUMA）优化和缓存数据对齐功能。

三、应用案例

（一）GBase 8a MPP Cluster 采用第五代英特尔至强可扩展处理器实现 1.2 倍性能提升，有效提高数据压缩性能

南大通用推出了大规模分布式并行数据库集群系统 GBase 8a MPP Cluster。为了提升 GBase 8a MPP Cluster 的压缩率以及性能，南大通用采用第五代英特尔®至强®可扩展处理器，以及处理器集成的高级硬件特性 — 英特尔®存内分析加速器（英特尔® IAA），在轻松实现超过 1:20 的压缩比的同时，能够在不同压缩算法中，实现高达 1.2 倍的性能增长。

南大通用测试了在第五代英特尔®至强®可扩展处理器中，启用英特尔®IAA 加速器前后的压缩率与性能表现。测试数据如图 4 所示，对比 ZSTD 压缩算法，英特尔® IAA 的性能提升了 10%；对比 Rapidz 算法，英特尔® IAA 压缩率提升了 33%。随后，南大通用还测试了在 GBase 8a MPP Cluster 数据压缩负载中，第四代/第五代英特尔®至强®可扩展处理器的代际性能差异。测试数据如图 5 所示，对比第四代英特尔®至强®可扩展处理器，第五代英特尔®至强®可扩展处理器对不同压缩算法的性能提升最多可达 20%。

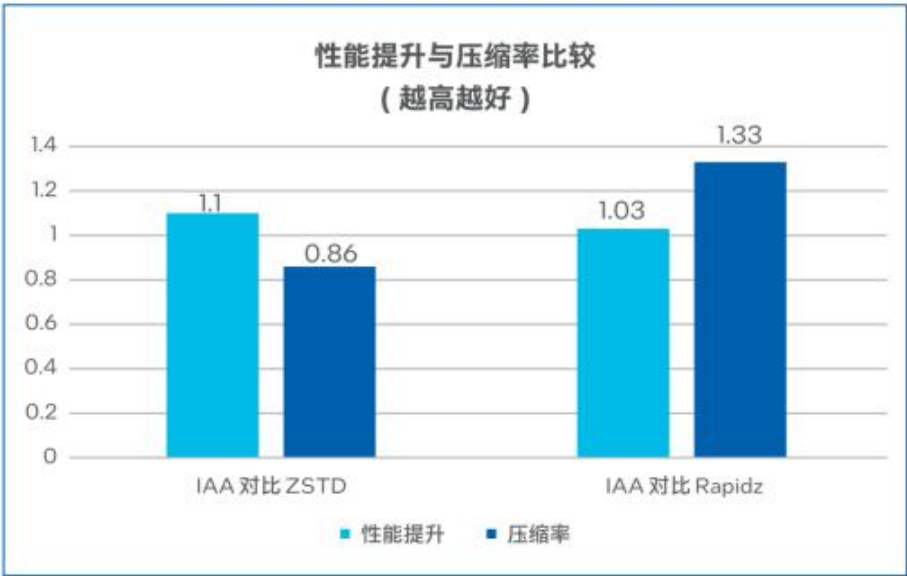


图 6：英特尔®IAA 启用前后的性能与压缩率差异

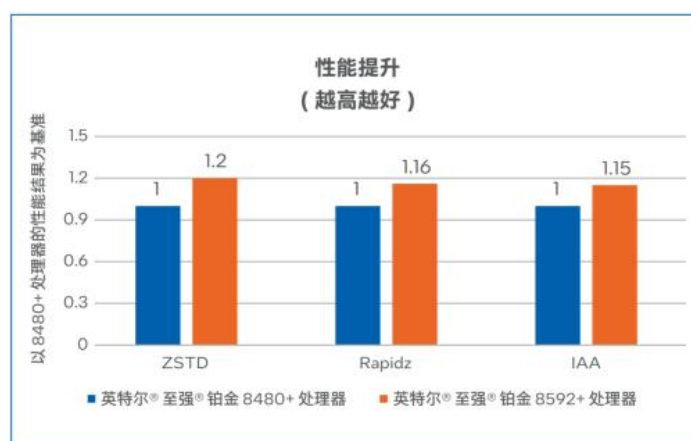


图 7：第四代/第五代英特尔®至强®可扩展处理器在不同压缩算法的代际性能差异

本节案例内容来自：<https://www.intel.cn/content/www/cn/zh/now/data-centric/gbase-8a-mpp-cluster-improve-data-compression.html>

（二）第五代英特尔®至强®可扩展处理器助力星环科技分布式向量数据库 Transwarp Hippo 实现 2 倍性能提升

星环科技分布式向量数据库 Transwarp Hippo 作为一款企业级云原生分布式向量数据库，基于分布式特性，可以对文档、图片、音视频等多源、海量数据转化后的多维向量进行统一存储和管理。它能够通过水平扩展架构，充分发挥并行检索能力，实现毫秒级高性能数据检索，结合相似度检索等技术，帮助用户快速挖掘数据价值。

为了进一步提升性能表现，星环科技验证了基于第五代英特尔®至强®可扩展处理器的分布式向量数据库 Transwarp Hippo 的性能表现。星环科技选用了 Transwarp KNN search 评测程序，该评测程序模拟用户的 top K 邻近范围查询。测试数据如图 6 所示，对比第三代英特尔®至强®可扩展处理器，基于第五代英特尔®至强®可扩展处理器的 Transwarp Hippo 性能是其 2.07 倍。

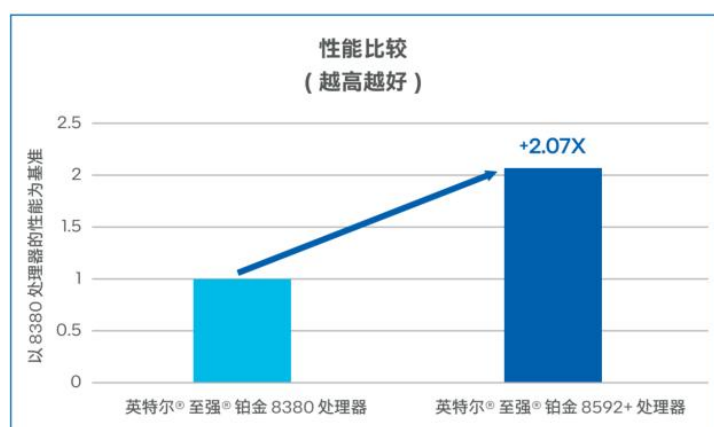


图 8：基于不同代际英特尔®至强®可扩展处理器的 Transwarp Hippo 性能对比

本节案例内容来自：<https://www.intel.cn/content/www/cn/zh/now/data-centric/distributed-vector-database-transwarp-hippo.html>

（三）第四代英特尔®至强®可扩展处理器助力 TiDB 开源分布式数据库大幅提升性能，优化存储空间

TiDB 是 PingCAP 公司自主设计、研发的开源分布式数据库，是一款同时支持在线事务处理与在线分析处理（Hybrid Transactional and Analytical Processing, HTAP）的融合型分布式数据库产品，具备水平扩容或者缩容、金融级高可用、实时 HTAP、云原生、兼容 MySQL 协议和 MySQL 生态等核心特性。TiDB 为用户提供一栈式联机事务处理过程（OLTP）、联机分析处理（OLAP）和 HTAP 解决方案，适用于高可用、强一致、数据规模较大等应用场景。

对于数据库应用而言，第四代英特尔®至强®可扩展处理器提供了更多的内核，以及更多的 Sub-NUMA Clustering (SNC) 节点，使得数据库系统能够实现明显的代际性能提升。

在 OLTP 场景中，为了验证 CPU 升级带来的性能提升，PingCAP 进行了测试，验证了在 Sysbench 基准测试中，英特尔®至强®铂金 8380/8480+处理器的只读、读写性能差异。测试数据如图 9、图 10 所示，基于英特尔®至强®铂金 8480+处理器的 TiDB 在 Sysbench 只读测试中性能达到基准配置的 1.62 倍，在 Sysbench 读写测试中性能达到后者的 1.43 倍。

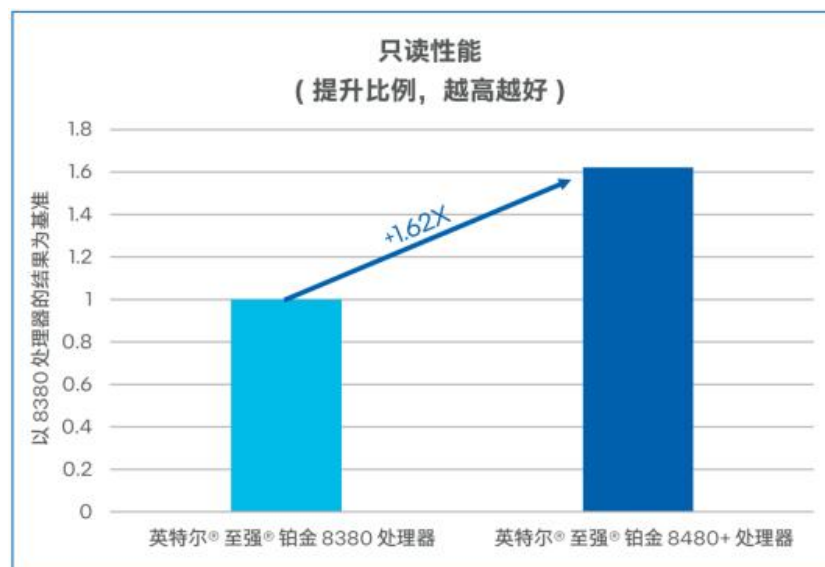


图 9：TiDB 只读测试性能

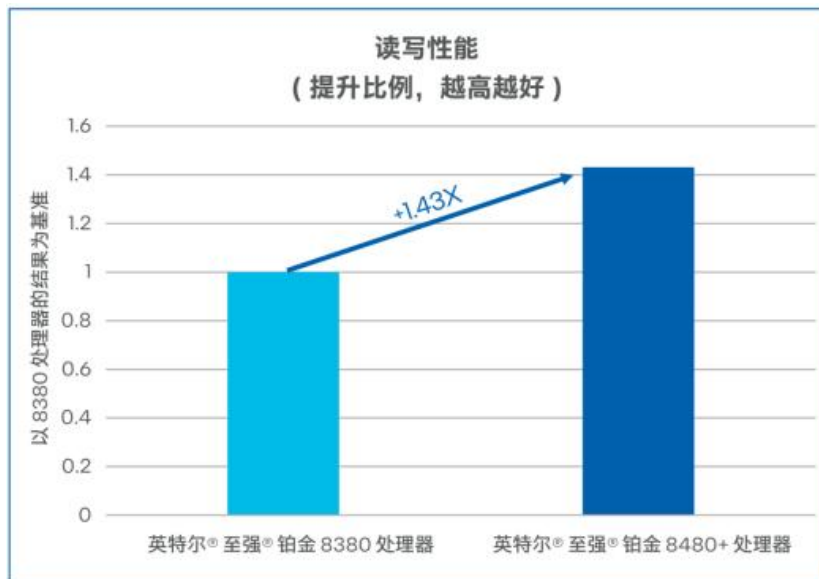


图:10: TiDB 读写测试性能

在 OLAP 场景中, PingCAP 还希望能够提升 TiDB 的海量数据压缩存储能力, 从而减少存储数据所需的空間, 降低相应的成本投入。为此, TiDB 采用了第四代英特尔® 至强®可扩展处理器集成的英特尔®IAA 加速器。

PingCAP 对比了在不同处理器配置下, 英特尔®IAA 以及 LZ4 无损压缩算法的压缩率差异。测试数据如图 11 所示, 采用英特尔®IAA 替代 LZ4 之后, TiDB 压缩率达到 LZ4 压缩算法的 1.4 倍, 主要针对列存储引擎 TiFlash 的使用场景, 能够大幅节省存储空间。此外, PingCAP 还测试了在不同的处理器与压缩算法的组合下, 数据库的性能差异。测试数据如图 12 所示, 在采用英特尔®IAA 替代 LZ4 进行压缩之后, 数据库的性能不仅没有降低, 还实现了一定的提升。

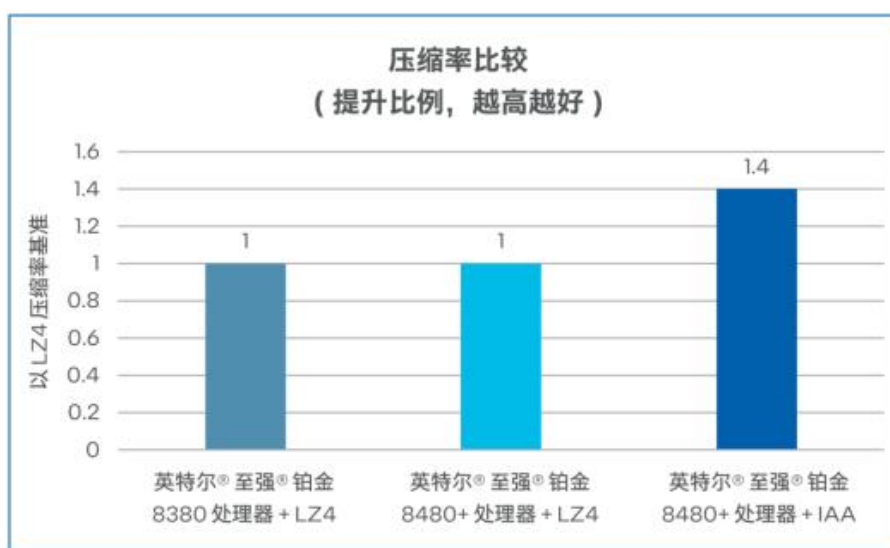


图:11: 不同处理器与压缩算法下的压缩率

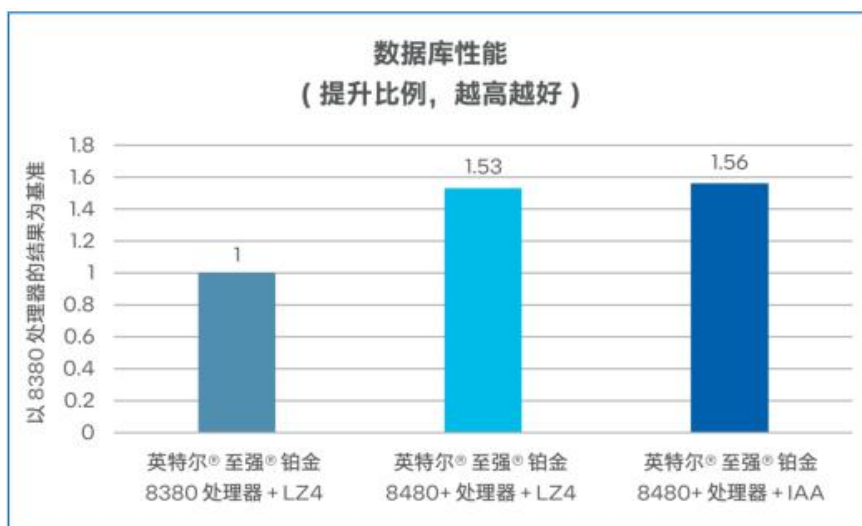


图 12：不同处理器与压缩算法下的性能差异

本节案例内容来自：<https://www.intel.cn/content/www/cn/zh/now/data-centric/boost-performance-optimize-storage-space-for-tidb.html>

（四）第五代英特尔® 至强® 可扩展处理器助腾讯云向量数据库成为大模型时代数据中枢

基于英特尔®AVX-512、英特尔®AMX、英特尔®FMAL 加速库等技术，腾讯云与英特尔一起，针对腾讯云向量数据库常用的一些计算库进行了专门的优化。包括：

1. **FAISS**: 方案中针对其不同的索引提出了不同的优化方案，包括面向 IVF-FLAT 算法的 ReadOnce（单次读取）和 Discretization（离散化）两种优化思路，以及借助英特尔®AVX-512 加速 IVF-PQFastScan 算法和 IVF-SQ 索引的优化方案；
2. **HNSWlib**: 方案借助英特尔®AVX-512, 对 HNSWlib 的向量检索性能进行了加速。同时方案也针对增删数据后的性能和召回率抖动的问题进行了专项优化，使 HNSWlib 的性能和召回率可以保持较平稳状态。

场景 1：英特尔®AVX-512 优化效果与代际性能提升测试

测试分为以下三个对比组：

- (1) 基准组：基于第三代至强®可扩展处理器的腾讯云 S6 服务器，实例规格：16 虚拟核；测试中使用 Faiss 计算库的 IVF-PQFastScan 算法进行检索，向量维度 768，数据集数据量 10 万，nprobe=10；
- (2) 测试组 1：基于第三代至强®可扩展处理器的腾讯云 S6 服务器，实例规格：16 虚拟核；使用英特尔®AVX-512 对 Faiss 计算库的 IVF-PQFastScan 算法进行优化，向量维度 768，数据集数据量 10 万，nprobe=10；

(3) 测试组 2：基于第五代至强®可扩展处理器的服务器，实例规格：16 虚拟核；使用英特尔®AVX-512 对 Faiss 计算库的 IVF-PQFastScan 算法进行优化，向量维度 768，数据集数据量 10 万，nprobe=10。

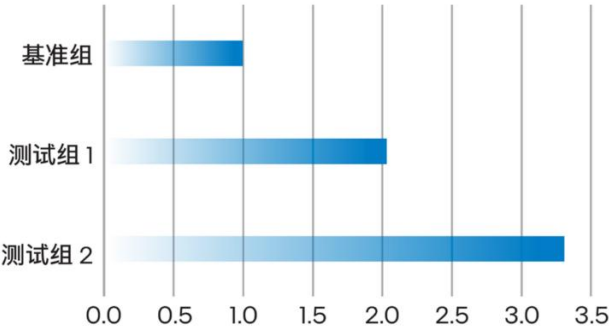


图 13：英特尔软硬件产品与技术带来的性能提升（归一化）

测试结果如图 13 所示，经数据归一化对比后，在同样使用腾讯云 S6 服务器（基于第三代至强® 可扩展处理器）的情况下，使用英特尔®AVX-512 优化后，使用 IVF-PQFastScan 算法执行向量检索时的 QPS 性能提升了约 100%，而将算力设备升级为第五代英特尔®至强®可扩展处理器后，性能相比基准组提升了约 230%。这表明，第五代英特尔®至强®可扩展处理器与英特尔®AVX-512 指令集能显著提升向量数据库的向量检索效率。

场景 2：英特尔® AMX 进一步提升向量检索性能

测试分为以下两个对比组：

- (1) 基准组：基于第五代至强®可扩展处理器的服务器，规格：16 虚拟核，使用经英特尔® AVX-512 优化的英特尔®FMAL 执行暴力搜索算法，数据量为 1,000 万，batch size 为 16，数据格式为 FP32，线程数为 16；
- (2) 测试组：基于第五代至强®可扩展处理器的服务器，规格：16 虚拟核，使用经英特尔®AMX 优化的英特尔®FMAL 执行暴力搜索算法，数据量为 1,000 万，batch size 为 16，数据格式为 INT8，block size 为 320，线程数为 16。

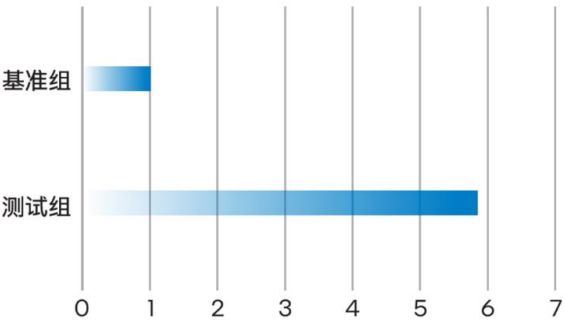


图 14：英特尔® AMX 优化加速暴力检索的吞吐性能（归一化）

测试结果如图 14 所示，经数据归一化对比后，在同样使用第五代英特尔® 至强® 可扩展处理器的算力平台上，使用英特尔®AMX 加速数据格式为 INT8 的测试场景相比使用英特尔®AVX-512 加速数据格式为 FP32 的测试场景，性能提升达约 5.8 倍。这表明，第五代英特尔® 至强® 可扩展处理器与英特尔®AMX 的配合，能进一步帮助腾讯云向量数据库提升向量检索效率。

本节案例内容来自：<https://www.intel.cn/content/www/cn/zh/customer-spotlight/cases/tencent-cloud-vector-db-data-hub-era-big-models.html>

（五）基于第四代英特尔®至强®可扩展处理器的星环科技 Transwarp Hippo 分布式向量数据库实现性能显著提升

星环科技 Transwarp Hippo 分布式向量数据库支持存储、索引以及管理海量的向量式数据集，提供向量相似度检索、高密度向量聚类能力，有效地解决了大模型知识时效性低、输入能力有限、准确度低等问题，让大模型更高效地存储和读取知识库，降低训练和推理成本，激发更多的 AI 应用场景。在赋予大模型拥有长期记忆的同时，还可以协助企业解决目前最担忧的大模型数据隐私泄露问题。产品架构如图 15 所示。

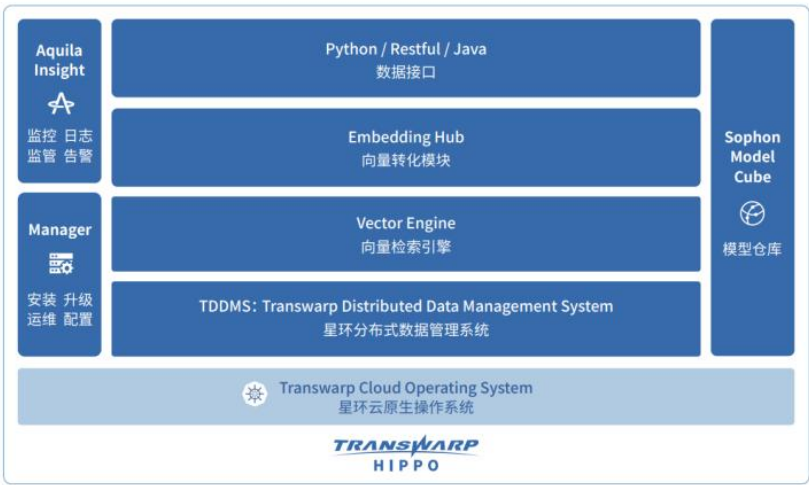


图 15：星环科技 Transwarp Hippo 分布式向量数据库架构

第四代英特尔®至强®可扩展处理器与星环科技 Transwarp Hippo 分布式向量数据库软硬件融合的深度优化：

1. 使用 AVX512 扩展指令集重写向量距离计算函数，显著降低向量计算需要的 CPU 指令数量与 CPU 时钟周期，充分发挥第四代英特尔® 至强® 可扩展处理器高内存带宽的优势。

- 2. NUMA 友好的向量计算负载调度算法，避免 CPU 远程内存访问造成 CPU 阻塞，充分发挥第四代英特尔® 至强® 可扩展处理器多核性能的优势。
- 3. 基于数据离散度的浮点数矢量化算法，充分利用 VNNI 指令集，进一步提升向量计算性能。

通过配置第四代英特尔® 至强® 可扩展处理器，星环科技 Transwarp Hippo 在向量索引层面实现了 20% ~ 30% 的性能提升，可全面满足个性化推荐、智能问答、大模型应用等场景对向量数据库系统计算能力的要求。

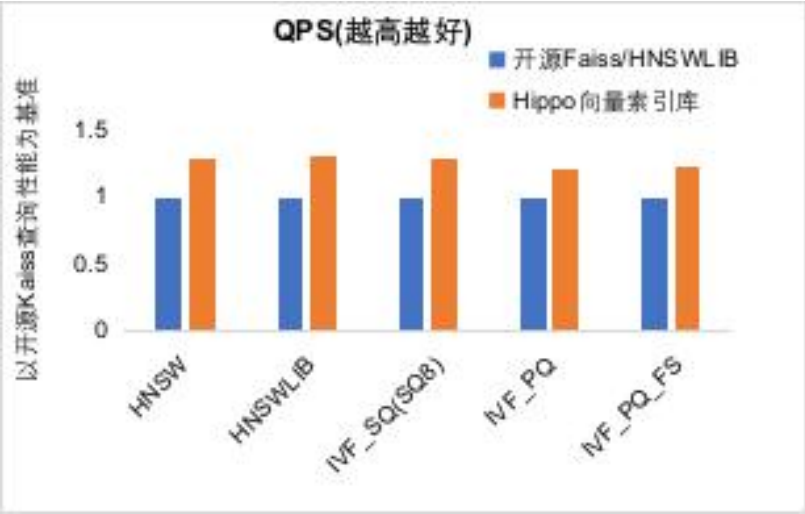


图:16: Hippo 向量索引库与开源 Faiss 查询性能对比

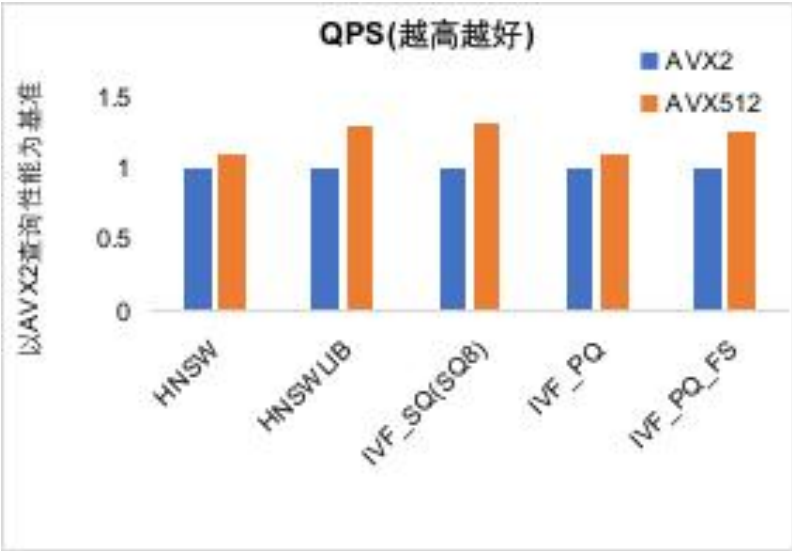


图 17: Hippo 向量索引库 AVX2/AVX-512 查询性能对比

本节案例内容来自：<https://www.intel.cn/content/www/cn/zh/artificial-intelligence/analytics/transwarp-distributed-vector-database-hippo-llm.html>

（六）英特尔 oneAPI 工具大幅提升腾讯云数据库 MySQL 的性能

腾讯实现了数据库托管服务腾讯云数据库 MySQL 性能的大幅提升，这一服务基于开源关系型数据库管理系统 MySQL，在英特尔® 至强® 可扩展处理器上开发而成。

利用英特尔® oneAPI DPC++/C++编译器，腾讯以结合链接时优化 (LTO) 和配置文件引导优化 (PGO) 的方法，构建了高性能 MySQL。通过链接时优化，编译器对应用程序进行模块间优化 (IPO)，允许对代码实现深入分析和进一步的优化，来达到更好的性能。配置文件引导优化则向编译器提供程序中最常被执行区域的信息。这些技术相结合，共同使腾讯云数据库 MySQL 的性能得到显著提升，最高可达 85%。

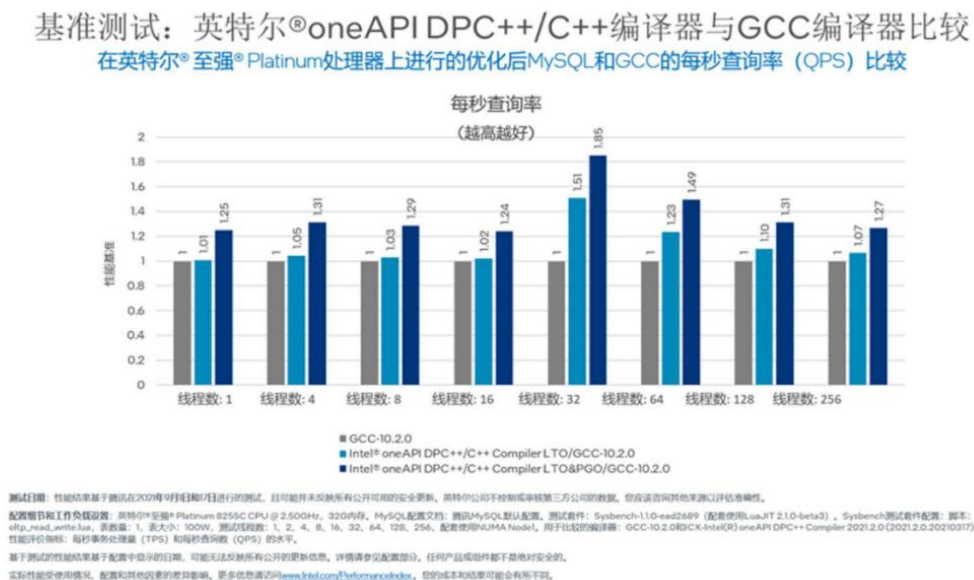


图 18：英特尔® OneAPI DPC++/C++编译器与 GCC 编译器 MySQL 性能比较

本节案例内容来自：<https://www.intel.cn/content/www/cn/zh/newsroom/news/oneapi-enhances-performance-of-tencentdb-for-mysql.html>

总结：

除了以上提到的案例，第四代/第五代英特尔®至强®可扩展处理器在更加广泛的数据库场景有着巨大潜力，有助于为数据库中密集的 IO 操作提供强大的算力支撑，同时在能耗比、总体拥有成本 (TCO) 方面有卓越的表现，可充分满足基于创新技术的数据库基础设施的构建需求，为用户提供敏捷、灵活、高性能、高可用的数据处理与分析方案。

2、存储

一、数据库存算分离/存算协同架构创新

（一）数据库存算分离前世今生，缘何成为产业发展重点

存算分离对于数据库从来不是新鲜话题。时光倒回到 20 世纪 80 年代，当时的核心业务大部分采用 IBM 大型机承载，大型机内融合数据库、计算单元、网络单元、存储单元等。很快，大型机的存储扩展性瓶颈逐步凸显，业务多元化也给数据共享提出更高诉求，企业级存储恰逢其时地推出来，最早是 EMC symmetrix，后来日立、惠普等厂商也推出自己的专业存储产品，专业存储逐步成为大型机的最佳伴侣。再后来，小型机的流行加速了 IT 解耦化，计算单元、存储单元、数据库单元被分拆出来，Oracle、Informix 等数据库开始流行，进入了数据库发展的黄金期。时至今日，在全球绝大部分数据库使用场景下，数据库+小型机/服务器+专业存储的存算分离架构都成为标配。

那么，为何今天数据库产业要重提存算分离技术？这源于中国数据库产业独特的发展路径。2010 年左右，一股“去 Oracle”的风潮从互联网行业刮起。由于互联网行业独特的超高并发业务压力和成本缩减诉求，一种采用开源数据库+PC 服务器+服务器本地盘的数据库建设架构在互联网厂商间流行开来，这种架构一般称为存算一体架构，因为存储和算力都是在同一个台服务器里面。由于开源数据库能力与 Oracle 在性能和可靠性上差距较大，一般都采用多个数据库节点并行工作承载过去一个 Oracle 数据库的工作，这就是分布式数据库。因此，分布式数据库在诞生之初是采用存算一体架构的。

客观来说，分布式数据库早期选择存算一体有其特殊的历史原因。2010 年前后，中国 IT 基础设施还处于发展初级阶段，从数据库到小型机、专业存储，都缺乏先进且可控的技术，这使得国内 IT 建设完全受制于国外厂商，不仅成本高昂，其能力也不符合中国国情，且服务支撑速度无法跟上互联网的演进速度。而存算一体+分布式架构，可以采用稳定性低、价格相对便宜的 PC 服务器和磁盘，建设一个可靠性满足互联网业务要求的数据库系统，且数据库直接访问本地硬盘路径更短、性能更高。因此存算一体架构的确是当时最满足互联网企业诉求的数据库部署方式。

今天，随着数字化与自主创新的诉求愈加强烈，在互联网率先成功的分布式数据库开始向更多行业扩展，并承担“替换 Oracle”的重任。但在金融、运营商等行业使用过程中，存算一体架构的分布式数据库逐步暴露出一些固有问题，使得产业界再次开始研讨分布式数据库是否应该走向存算分离架构。

（二）存算分离解决固有难题，逐步成为数据库建设标准

要谈为何存算分离，首先需要看看存算一体有何问题。存算一体架构存在的问题主要有两个，一个是数据失去共享能力，二是数据存储面对的所有问题都需要 DBA 去解决：

1. **数据共享问题**：由于服务器本地盘只能被本服务器使用，数据不具备跨服务器共享访问的能力，这使得服务器故障=数据丢失，为了提升数据库可用性，不得不把一个数据存多个副本。这引入了两个问题，一是空间的浪费问题，多副本+RAID 下平均硬盘利用率大都不到 10%；二是副本间的数据同步问题，副本过多时完成同步非常影响数据库性能。此外，分布式数据库往往会对数据库进行分库分表和数据分片，不同分片日积月累后会出现数据分布不均，此时就会出现部分硬盘闲置、部分硬盘爆满的情况，但闲置的资源并不能分配给资源不足的硬盘使用。

2. **存储功能问题**：事实上，数据存储并不仅仅是把数据写到盘里那么简单，还有很多问题需要解决。比如说，当数据从 CPU 写入盘时发生了跳变，如何及时发现并纠正；当硬盘因为寿命或故障原因出现慢盘，如何及时隔离避免影响上层业务；此外还有诸如 RAID、重构、数据备份与恢复、数据加密等等复杂的数据存储功能。这些功能在行业中有大量实在的需求，也有非常悠久的历史。然而如果采用存算一体架构，由于服务器本地盘不具备以上所有能力，这些能力都需要 DBA 在应用层重新构建，其难度可想而知，且对业务影响也非常大。

随着存算一体的问题逐步明朗，存算分离再次获得广泛讨论。归根结底，过去分布式数据库是由专业技术人员和运维人员极其丰富的互联网大厂开发和维护的，所有问题可以及时内部闭环，业务中断的后果也并不那么严重；但传统行业，缺乏技术和运维人员，业务重要性却极高，这一矛盾使得存算一体分布式数据库无法匹配传统行业 DBA 们的诉求；除了数据库以外，在其他 IT 环节他们需要“懂行人”。目光再次转向专业 IT 基础设施。在数据库历史中，专业存储作为数据存储“懂行人”和数据库最佳伴侣，已在传统行业扎根多年，早已具备了以上两大问题的完美答案。只需要将分布式数据库与之简单对接，所有问题都能迎刃而解。简单来说，就是“让专业的人，干专业的事情”，把数据存储的所有麻烦问题都交给专业存储去做，DBA 只需要关心上层的业务逻辑即可。具体来说，在一套专业存储构建的存储池当中，划分多个不同的虚拟存储空间（LUN 或文件系统），分配给不同服务器使用，而数据库无需修改，就像使用本地盘的方式一样使用专业存储映射过来的存储空间即可。

使用存算分离架构后，数据库立马可以具备各种专业级存储能力，再也不用为数据存储而头疼。举个例子，过去分布式数据库最担心出现慢盘，导致数据访问迟迟无响应，业务出现宕机；而通过专业存储的慢盘隔离，即使出现慢盘也会被存储及时隔离，不会

影响业务，且可及时通知运维检查更换。再比如，分布式数据库分片不均，过去会出现资源浪费，而在存算分离架构下，通过共享和 Thin LUN 技术，各个节点按需索取存储空间，硬盘资源 0 浪费。存算分离其他优势此处不再一一赘述，这些优势使得分布式数据库可以很快满足非互联网行业客户诉求，因此正逐步成为新的建设标准。

此外，在 2010 年前后采用“存算一体”架构的历史原因也不复存在了。如今，从服务器到交换机再到存储，中国品牌都已取得巨大进步，尤其是专业存储从技术到收入均已达到世界顶尖的水平，获得了来自全世界客户的认可。中国企业再也不用担心用不到最好的产品和最及时的服务。而 NVMe 及 NVMe over Fabric 技术的出现，也大大拉近了本地盘和外置存储的性能差距，如今时延上二者已接近同等水平，而并发吞吐上外置存储大幅领先本地盘。这无疑将加速数据库存算分离标准建设。

（三）从“存算分离”走向“存算协同”，数据库迈上新台阶

在数据库迈向存算分离以后，数据库产业开始思考下一步问题：专业存储和数据库能否产生更多的合作点？这一问题的产生有两个背景，一是随着数字化的推进，IT 业务压力越来越大，而且出现了如 AI 这类高密度数据访问应用，算力成为稀缺资源，产业开始寻求通过算力分摊的方式来解决部分业务诉求；二是随着新型硬件能力大幅提升，充分利用新兴存储硬件提升数据库性能，向着更高性能、更强扩展和更低总体成本的数据库发展成为产业发展的一个热门方向。一言以蔽之，“存算协同”的本质是“以存强算”，利用存储实现数据库所不能及或难以企及的能力。

数据库存算协同有很多方向。一个非常经典的例子是，过去 Oracle Exadata 中，可以将部分 SQL 处理或其他任务从服务器卸载到存储上去做，以实现近数据计算，并节省带宽资源；如今分布式数据库也可以在这方面发力。以下仅举两个存算协同案例供产业同仁借鉴：

1. **存储同步复制助力数据库容灾**：从 Oracle ADG 开始，数据库容灾大都在软件层实现，通过日志流式复制来实现数据库容灾。这一方案存在的问题是，生产库/主库需要等待日志完全写入从库时才能提交事务，因而性能会受到较大影响；而且当从库与主库在不同数据中心时，复制网络的抖动、阻塞将造成主库性能受到严重影响甚至宕机。专业存储具有非常强大的数据复制能力，且具有非常强大的复制误码检测、抖动自动检测/切换机制，通过将日志复制下沉到存储，可以避免主库等待从库写入造成性能损失甚至宕机；甚至可以在容灾数据中心设置一套容灾数据库集群，通过实时回放日志，达到双集群容灾 RPO=0 的效果，可以达到金融核心级的容灾要求。这一技术已在 GaussDB/Vastbase 等数据库上实现落地。

2. **存储文件系统助力数据库多写多读**：Oracle RAC 是数据库多写多读架构的标杆，其中最难实现的技术，一个是缓存层的池化，一个是存储层的高效共享。Oracle 构筑了一个共享能力极强的文件管理系统 ASM，但对 PC 服务器的能力要求很高；如今这一功能可以在存储上完成。企业级 NAS 存储已有多年的积累，在多节点共享访问能力和企业级功能上已完全具备，但其在性能与可靠性上还达不到数据库要求。如今一项新的文件系统技术 NFS+ 应运而生，它可以将过去文件系统客户端并发访问端口数从 1 扩展到 4，并具备链路故障快速切换的能力，使文件系统从性能到可靠性上都满足数据库要求。通过数据共享的下沉，服务器资源被释放并更高效地执行 SQL 解析等工作，可以高效支撑 Cantian 引擎等数据库多写多读使能引擎，且在高并发场景下性能表现更优。

加速数据库与存储间的“存算协同”，需要数据库厂商在自我研发、创新同时，进一步加强与存储厂商的兼容性适配及联合创新力度，充分利用中国存储在核心应用复杂、高压、苛刻场景下已经取得的产业优势和成熟经验，共同对数据库方案进行打磨和优化，补齐在核心技术、灾备能力、技术服务能力的短板，使其能更好地服务于中国千行百业客户。在产业与行业牵头单位的统筹规划下，我国将充分发挥体制优势，加速“存算协同”发展，推动中国数据库从跟随快速迈向领先。

二、可计算存储助推算力时代下的数据库性能

（一）可计算存储的诞生背景

1. 什么是可计算存储

根据 SNIA(全球网络存储工业协会)的定义：可计算存储(Computational Storage)被定义为一种提供将计算功能与存储耦合以卸载(Offload)主机处理或减少数据搬移的架构，该架构通过将主机与存储驱动器之间的计算资源进行集成，或直接与存储驱动器内的计算资源进行集成，以提高应用程序性能或优化基础设施的效率，目标是实现并行计算，或(和)减轻对现有的计算、内存、存储和 I/O 的限制。

可计算存储驱动器(Computing Storage Drive, 缩写 CSD, 下文统称为 CSD)是一种符合可计算存储架构定义的驱动器实现形态。通俗地讲，CSD 可以简单地理解为是一种在传统 NVMe SSD 基础上进行迭代，且符合可计算存储架构定义的新型 NVMe SSD，CSD 和传统 NVMe SSD 的本质区别在于，CSD 在盘内(主控芯片内)集成了计算加速引擎，能够做到对应用透明无感知(无需对应用做任何改动)的前提下，实现对数据写入和读取操作的加速，从而实现应用性能加速的效果。

2. 存储面临的挑战

当今我们正处在一个前所未有的数字化转型的进程中，各种新兴技术的产生和使用都会面临着一个共同的问题，那就是数据产生和使用呈爆发性增长，这会给底层的计算和存储技术带来巨大的挑战。

在过去的几十年中，存储技术从卡带到磁盘，再到固态硬盘，容量和性能都得到了巨大的提升，但其提升的速度远远赶不上数据增长的速度。如果我们把 2020 年全球存储的产能加起来，大约 20ZB（相当于 20 亿张 10TB 的硬盘），这已经是比较惊人的产量，但到了 2025 年，根据国际数据公司(IDC)的预测，全球数据量将会达到 175ZB，但与此同时，存储的产能只能达到 22ZB，可想而知这将是存储面临的巨大挑战。

3.算力面临的挑战

半导体工业经历了将近半个世纪的高速发展，CPU 的性能每隔 18 个月翻一倍，价格下降一半。但是在最近的 10 多年里，由于半导体工艺逐渐接近物理极限，CPU 的性能每隔 18 个月的提升已经不足 2 倍，但与此同时数据的增长量却呈爆发式增长，这种情况下算力也将面临巨大的挑战。

当传统的计算与存储方式难以满足数据增长需求的时候，就必须通过创新来解决计算和存储的效率。要提升计算和存储的效率，最有效的解决方案就是将计算分流，例如：

（1）将高密度 AI 计算从 CPU 分流到 GPU/TPU 等这些对 AI 数据计算效率更高的专有芯片；

（2）将与网络相关的数据处理从 CPU 分流到智能网卡芯片；

（3）将作用于海量数据上的操作（例如：数据压缩与解压、数据过滤、数据加密与解密等）从 CPU 分流到存储设备内的主控芯片。

在上述计算分流方案中，将存储的数据处理分流到存储设备内部，理论上不但能节省主机 CPU 算力资源，而且不需要占用系统总线来回拷贝数据。随着存储容量的扩展，还能线性扩展存储内处理数据的计算能力，甚至可以利用多块盘内的专有芯片并行处理更多的计算任务。

如此，催生带计算能力的存储设备势在必行，CSD 便应运而生。

（二）可计算存储的应用探索

可计算存储的基本原理并不复杂，下面与传统存储架构进行简单的比较说明。

1.传统存储

数据从存储设备移动到主机内存，由主机 CPU 计算处理完成之后，再经由主机内存回到存储设备内。当数据量不断增长，或数据的处理需求不断增长时，有效的解决方案就是扩展硬件（如：更大的存储、更快的 CPU、更多的服务器等）。但扩展硬件的速度往往赶不上数据增长的速度，更快的数据增长意味着存储需要存放更多的数据、有更多的数据需要从存储设备移动到主机内存、CPU 需要处理更多的数据，这个时候 CPU 算力和存储容量这两侧以及系统总线带宽资源都可能成为瓶颈，所以扩展硬件的方案在这种情况下就显得捉襟见肘了。

2.可计算存储

拉近数据存储端（存储设备）与数据运算处理端（专有芯片）的距离，直接利用盘内集成的专有芯片在盘内完成数据的处理运算。如此一来，不但可大幅减少数据的搬移量，还可以大幅提升计算的性能（专有芯片的计算性能远远高于通用芯片）。如果在一台服务器中插有多块可计算存储设备，甚至还可以利用多块可计算存储设备中的专有芯片来大幅提升计算的并行度。

在可计算存储领域，来自 ScaleFlux 的 CSD 系列 NVMe SSD（下文简称 CSD）就是这样一款自带专有芯片与计算引擎的可计算存储产品。为便于大家更好地了解这款产品，下面我们对其进行简单地展开介绍：

（1）在接口形态上，CSD 支持 U.2 和 AIC（半高半长）两种标准的物理接口，支持 NVMe 协议；

（2）如果把 CSD 的盘拆开来看（以 U.2 接口形态的产品为例，如图 1 所示），CSD 其实与普通 NVMe SSD 并无多大区别，都包含存储芯片（闪存颗粒）和主控芯片，但不一样的地方在于可计算存储的主控芯片内包含了计算加速引擎以及配套的核心软件与固件，可以更好地发挥底层硬件的性能，以便更好地服务上层核心应用（如数据库应用等）。

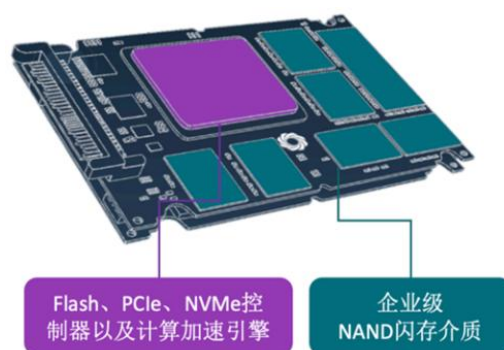


图 1：CSD 的内部架构

那么，对于完成了 CPU 计算任务卸载的同时还能够大幅提升应用性能的“隐形之手”——计算加速引擎，它究竟是一种什么黑科技呢？

在 CSD 这款产品中，可圈可点的计算加速引擎当属“透明压缩/解压”。新技术的应用探索过程通常都少不了蜿蜒曲折，可计算存储的应用探索过程也是如此。下面将对这种计算加速引擎的应用探索过程进行简单的介绍。

3.透明压缩/解压

这是利用可计算存储缓解容量压力并提升应用性能的探索尝试。

SSD 的出现极大地提升了存储性能（如提升 IOPS 并降低时延），成本也是不断下降。但是逐年下降的价格依旧跟不上数据量的爆炸式增长。SSD 自身的特性决定了其容量大小不仅会影响成本，也会影响性能。

对于空间复用（这里指删除旧数据之后，在腾出的空间中写入新数据），SSD 并不能像内存和传统机械硬盘直接覆盖旧数据。在使用 SSD 的场景中，当上层应用删除了某些数据之后，向旧数据占用的空间写入新的数据时，SSD 内部需要先对旧数据占用的闪存块执行擦除操作，之后在“干净”的闪存页中写入新数据。

在执行擦除操作的闪存块中，如果全部是旧数据（无效数据），则可以直接执行擦除，但在真实的应用场景中同一个闪存块内除了旧数据之外，也有可能存在正常文件的数据（有效数据），这个时候并不能直接擦除。在执行擦除操作之前，必须将这些有效数据“搬移”到其他空闲的闪存块中，然后才能执行擦除操作，这个过程叫垃圾回收。在这个过程中可能会产生写放大(Write Amplification)。

一旦 SSD 剩余空间显著变少，且出现大量数据碎片时，可能会频繁地触发垃圾回收，这过程中会产生严重的写放大。严重的写放大会严重影响 SSD 性能（如显著感知到 IO 延迟变大、IOPS 降低）。对于 SSD 内部的 IO 延迟（非应用感知的 IO 延迟），单个擦除操作的延迟是写操作延迟的几倍，而写操作的延迟又是读操作的几十倍。在混合读写的场景，垃圾回收会引发延时抖动，进而影响应用性能。

为降低垃圾回收的频率，SSD 不仅会优化垃圾回收的算法，同时也会预留一部分空间专门用于垃圾回收过程中的数据腾挪(预留空间通常叫 OP 空间: Over Provision)。对于 OP 空间的设置，企业级 SSD 的 OP 通常设置为 28%，消费级 SSD 的 OP 通常设置为 7%。SSD 内部空闲的空间越多，写放大就越小，而写放大减小可显著提升 SSD 的性能并降低 IO 延迟。

如果在盘内可直接对数据进行压缩，即可在不改变应用架构的前提下立竿见影地增加 SSD 内部空闲空间，不仅能节省存储设备的空间占用，更能优化存储设备的性能。那么，CSD 具体是如何做到的呢，下面我们稍微做一下展开说明：

(1) 如图 2 所示，相比不带硬件透明压缩的普通 NVMe SSD，内置了硬件透明压缩的 CSD，在对数据进行读写访问的 I/O 路径上，对数据的压缩与解压操作是在数据经过集成压缩引擎的主控处理器时，顺带就完成了对数据的压缩与解压。在进行数据写访问时，对数据的压缩时机是在主机端写入盘内之后，写入 NAND 闪存介质之前压缩的；在进行数据的读访问时，对数据的解压时机是从 NAND 闪存介质读取之后，返回给主机之前解压的，即后端 NAND 带宽中传输的数据始终是经过压缩后的数据，这在数据具备一定可压性的场景中，可显著降低后端 NAND 带宽的占用，从而达到降低应用 I/O 访问延迟，提升 IOPS 和吞吐的目的。在较重的访问负载场景中(如应用的并行访问会话数量大于 16 个时)，与普通 NVMe SSD 相比，CSD 可显著提升应用性能。

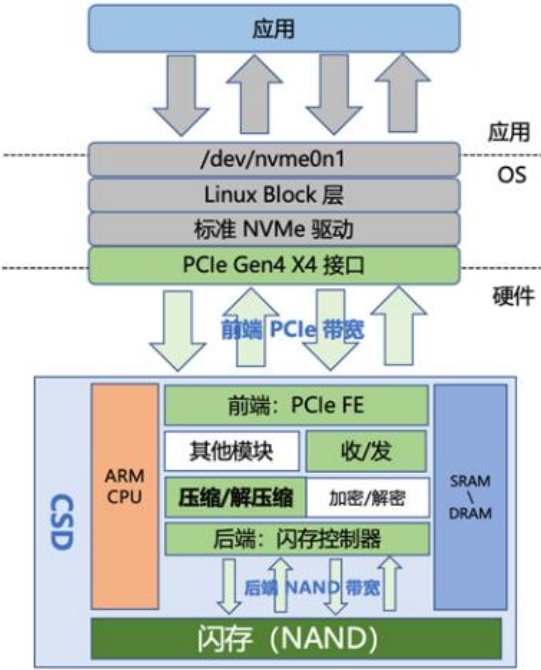


图2: CSD 降低后端 NAND 带宽占用原理示意图

(2) 如图 3 所示，相比不带硬件透明压缩的普通 NVMe SSD，内置了硬件透明压缩的 CSD，在数据具备一定可压性的场景中，可显著降低后端 NAND 闪存介质的空间占用。如下图所示，假设数据具有 2:1 可压性，在 CSD 中实际写入 NAND 闪存介质之前，在盘内主控处理器中经过压缩处理之后数据量已经减少了 50%，这就意味着与普通 NVMe SSD 在写入相同数据量的情况下，CSD 可减少 50% 的 NAND 闪存介质空间占用，可以简单理解为变相地增加了盘内的 OP 空间，前面我们也花了大量篇幅讲解 OP 空间的作用，更大的 OP 空间意味着可显著减少垃圾回收、数据搬迁的操作，也可显著减少写放大，更能显著提升存储设备的使用寿命，从而达到降低应用 I/O 访问延迟，提升 IOPS 和吞吐的目的。在数据量较大的场景中（如应用数据总量超过了存储设备可用容量的 50% 以上时），与普通 NVMe SSD 相比，CSD 可显著提升应用性能。



图 3：CSD 降低 NAND Flash 物理空间占用原理示意图

那么，既然压缩带来的好处显而易见，除了用硬件压缩之外，软件压缩能否达到类似的目的呢？显然也是可以的，目前对于市面上大多数的应用程序而言，或多或少都自带了一些压缩算法，对于压缩算法，目前业界提供了丰富多样的成熟算法，比如 zstd、zlib、brotli、quicklz、lzo1x、lz4、lzf、snappy 等。基于这些压缩算法，现行的压缩方案可简单归纳成软压缩（即消耗主机 CPU 的方案）和硬压缩（这里指可计算存储压缩方案）两种。为便于大家理解两者的区别，接下来我们对这两种主要的解决方案进行展开介绍。

（1）在软压缩方案（架构如图 4 所示）中，其压缩和解压能力完全由 CPU 提供算力。本质上是以“牺牲”CPU 资源来换取存储空间，因此，该方案存在如下几个突出的问题：

- ①CPU 抢占：会占用大量 CPU 资源，同时也会跟应用抢占 CPU 资源；
- ②数据复制导致的带宽抢占：在主存和 CPU 之间引入频繁且大量的数据复制（DRAM $\leftrightarrow$ L3 Cache $\leftrightarrow$ L2 Cache $\leftrightarrow$ L1 Cache $\leftrightarrow$ Registers），抢占服务器 PCIe 带宽和内存带宽，同时带来潜在的 CPU Cache Miss，进一步影响计算效率；
- ③时延不稳定：因为 CPU 抢占和带宽抢占，当操作系统负载较高时，操作系统中的时钟中断和任务调度增加了延迟的不确定性，这是 IO 密集型业务很难忍受的。

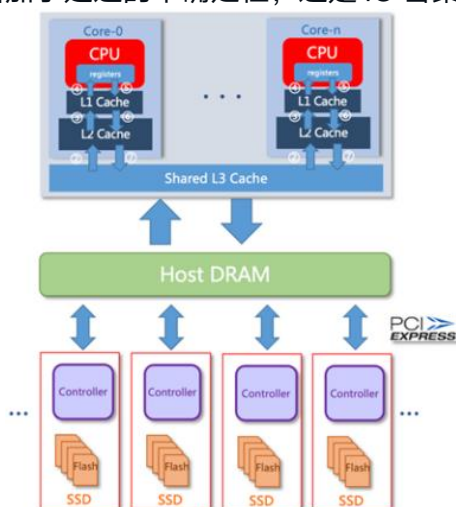


图 4：软压缩方案架构示意图

(2) 在硬压缩方案（架构图如图 5 所示）中，能够很好地解决软压缩方案中的突出问题：

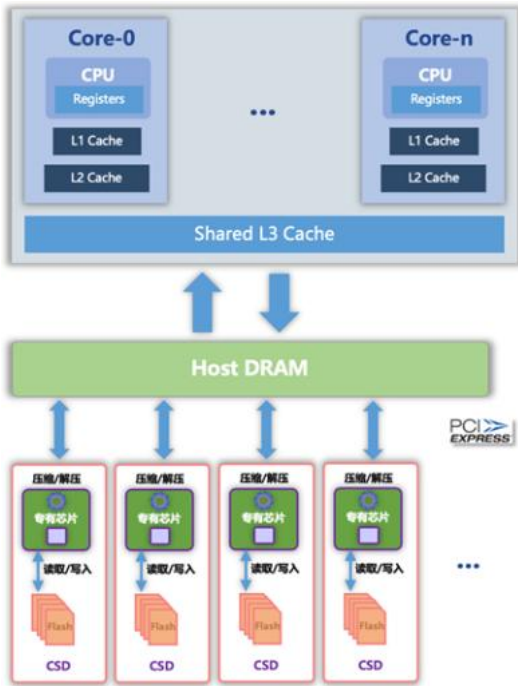


图 5：硬压缩方案(透明压缩)架构示意图

①CPU 负荷卸载：采用内置的专用芯片完成压缩和解压缩计算，实现 CPU 负荷卸载。专用芯片在低延时上具备天然的优势，非常适合计算密集型任务（比如矩阵运算、压缩和非对称加密）。首先，片上集成缓存和 DRAM 接口，减少与 CPU 交互，免于操作系统的进程调度和进程间干扰，从而提供可预测的时延。同时，专用芯片基于定制流水线 MIMD 设计，同时拥有流水线并行和数据并行，进一步降低时延。

②零拷贝：内置专用芯片的压缩和解压缩方式，不改变原有的数据路径，完全在盘内进行压缩解压任务，避免额外数据复制的同时，也大大降低了 IO 的延迟，这也是为什么可计算存储又叫“近”存储计算的原因；

③可线性扩展：每个可计算存储设备都内置压缩/解压引擎，在扩容存储容量的同时也能够线性扩展压缩/解压能力。

④压缩效率更高：一方面是因为采用硬件压缩可提升压缩效率，另外一方面是因为盘内压缩可将应用的所有类型的数据文件，以及文件系统 4k 对齐浪费的空间整体做压缩（软压缩方案通常只针对应用特定的索引文件或数据文件进行压缩，且因为文件系统的 4K 对齐是在应用的下层，软压缩方案无法对其进行压缩）。

综上所述，采用 CSD 硬压缩方案，数据压缩与解压的整个过程应用是毫无感知的，相比使用普通 NVMe SSD，应用在数据的读写操作上没有任何区别。当应用写数据时，

数据在写入盘之前保持着未压缩状态，当数据写入盘之后，通过盘内置的压缩引擎，结合内置的专有芯片对数据执行压缩，完成压缩之后再写入 NAND 闪存颗粒进行持久化；当应用读数据时，先从 NAND 闪存颗粒读出数据，然后通过内置的解压引擎，结合内置的专有芯片对数据执行解压，完成解压之后再返回给上层应用，不仅能够做到对应用透明，压缩效率也会更高，对应用提速的效果也更强，更能提升 SSD 设备的使用寿命，且基于压缩引擎实现的逻辑容量扩充更能，在数据具有较高可压性的场景中，更能显著降低整体基础设施的 TCO。

（三）可计算存储的应用案例

上文中我们花了很大篇幅讲解可计算存储以及相关的底层逻辑，那么可计算存储在应用场景中表现到底如何呢？

1.CSD 与普通 NVMe SSD 在 JESD219 模型下 FIO 基准测试工具性能对比：

（1）通过 FIO 基准测试工具(直接压测裸设备)，与普通 NVMe SSD 在同等场景下进行对比测试，结果表明，在应用数据具有较高压缩率的场景中，不但能够节省存储空间,还能够大幅提升 IOPS 与吞吐性能、降低 I/O 访问延迟,甚至还能够大幅提升 SSD 设备的使用寿命。

（2）如图 6 所示，在数据具有 2:1 以上压缩率时（意味着能够节省 50%以上的存储空间），CSD TLC SSD 在透明压缩/解压引擎的加持下，与同级别的普通 TLC SSD 相比，CSD 的 FIO 的随机写负载性能提高了 177%以上，FIO 的随机读写混合负载性能提高了 98%以上。

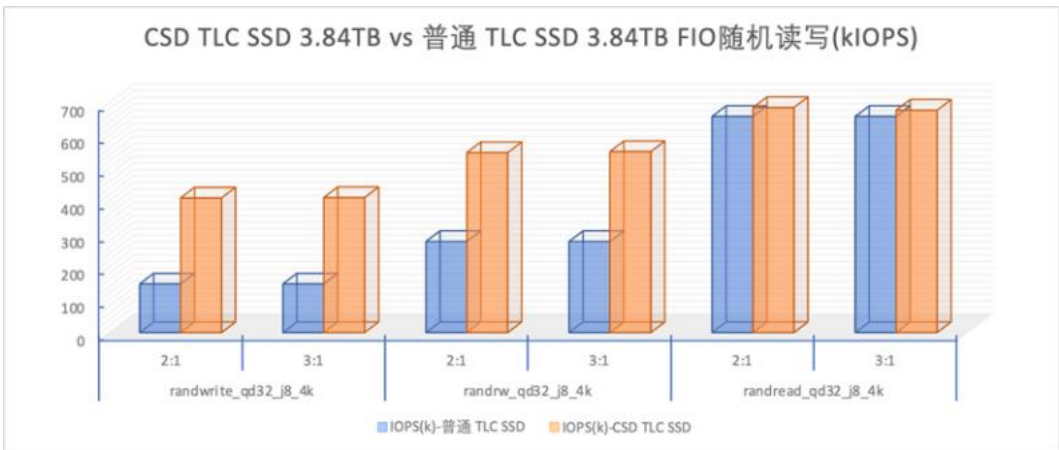


图 6：CSD 与普通 NVMe SSD 的 FIO 性能对比

2.CSD 在非常流行的 MySQL 与 PostgreSQL 两大开源关系型数据库中的性能表现：

(1) 在 Sysbench For MySQL 的基准测试中（如图 7 所示，样本数据压缩率 3:1），与同级别的普通 NVMe SSD 相比，CSD 在 Sysbench 的多种事务模型（图中的横轴坐标）、不同并发线程（图中的横轴坐标）的大部分负载场景中性能都远高于普通 NVMe SSD。且随着并发线程的逐步增加，CSD 的性能优势会更加明显（图中的纵轴表示性能提升比例）。在透明压缩引擎与盘内原子写特性（CSD 内置的特性，启用原子写时可关闭 MySQL 的双写缓冲功能）的加持下，CSD 的性能最高可提升 3 倍以上。

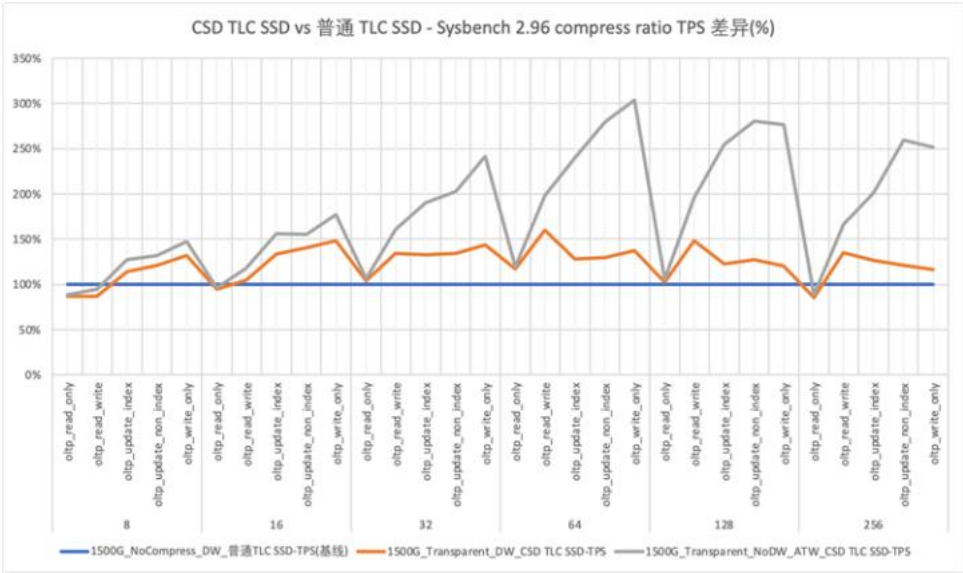


图 7： CSD 与普通 NVMe SSD 的 Sysbench For MySQL 性能对比

(2) 在 Sysbench For PostgreSQL 的基准测试中（如图 8 所示，样本数据压缩率 3:1），与同级别的普通 NVMe SSD 相比，CSD 在 Sysbench 的多种事务模型（图中的下横轴坐标）、不同并发线程（图中的上横轴坐标）的大部分负载场景中性能同样明显高于普通 NVMe SSD。在透明压缩引擎的加持下，CSD 的性能最高可提升 2 倍以上

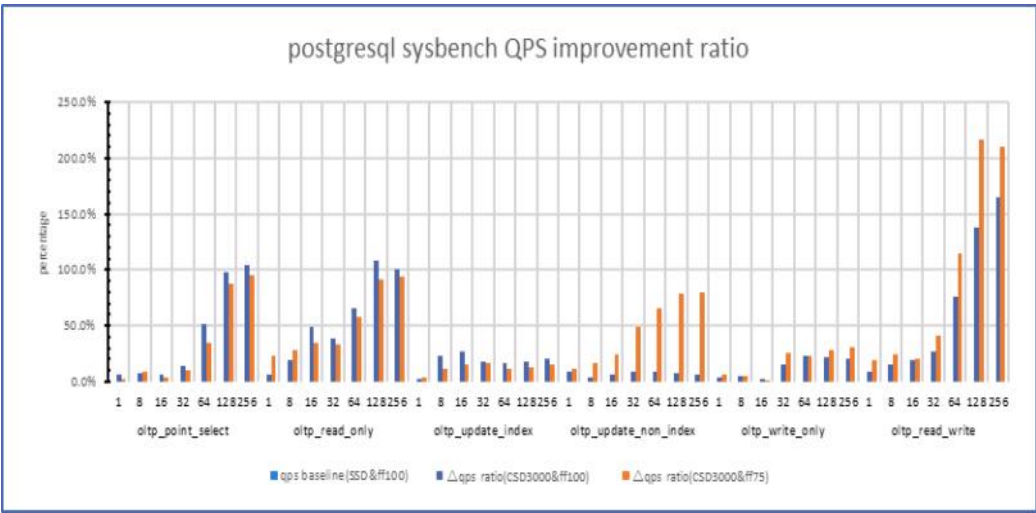


图 8： CSD 与普通 NVMe SSD 的 Sysbench For PostgreSQL 性能对比

3.CSD 在国产数据库 OpenGauss、达梦数据库、TDSQL 中的性能表现：

（1）在 Sysbench For OpenGauss 基准测试中（如图 9 所示，样本数据压缩率 3:1），CSD 的整盘压缩效率比 Sysbench For MySQL 基准测试中的表现更高，同时，在高并发的读写混合场景下与同级别的普通 NVMe SSD 相比,CSD 的性能最高可提升 60%以上。

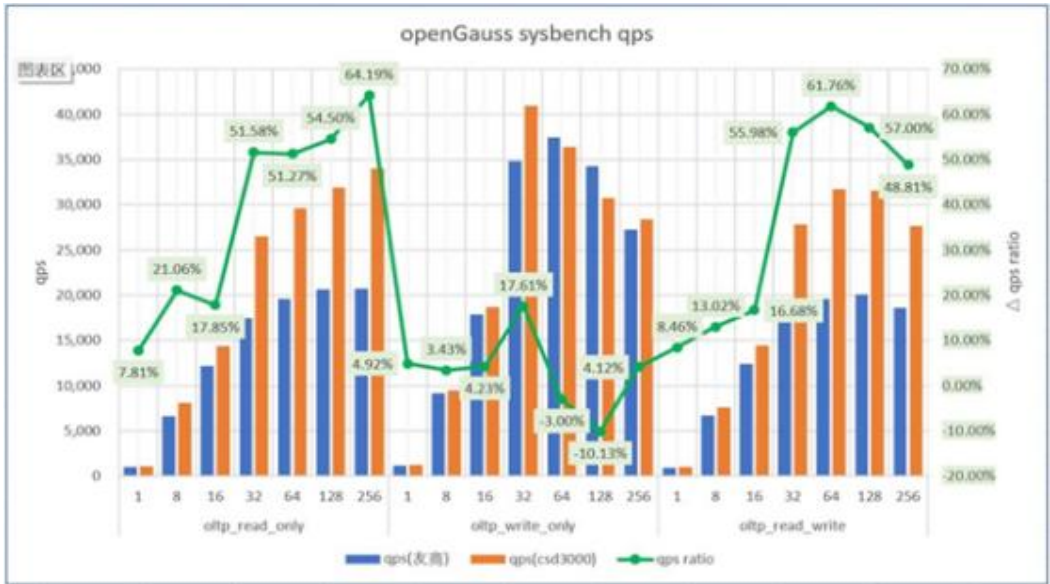


图 9：CSD 与普通 NVMe SSD 的 Sysbench For OpenGauss 性能对比

（2）在 Sysbench For 达梦数据库基准测试中（如图 10 所示，样本数据压缩率 3:1），与同级别的普通 NVMe SSD 相比，CSD 在不同的读写负载、不同的并发线程场景下，性能最高可提升 80%以上。

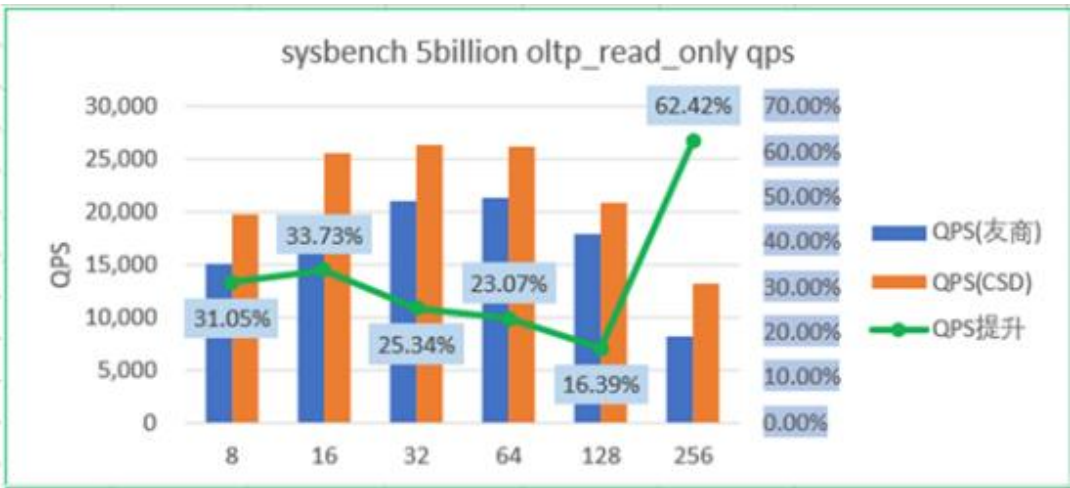




图 10： CSD 与普通 NVMe SSD 的 Sysbench For 达梦数据库性能对比

(3) 在 Sysbench For TDSQL 基准测试中 (如图 11 所示, 样本数据压缩率 2:1), 与同级别的普通 NVMe SSD 相比, CSD 在不同的读写负载、不同的并发线程场景下, 性能最高可提升 1.5 倍以上

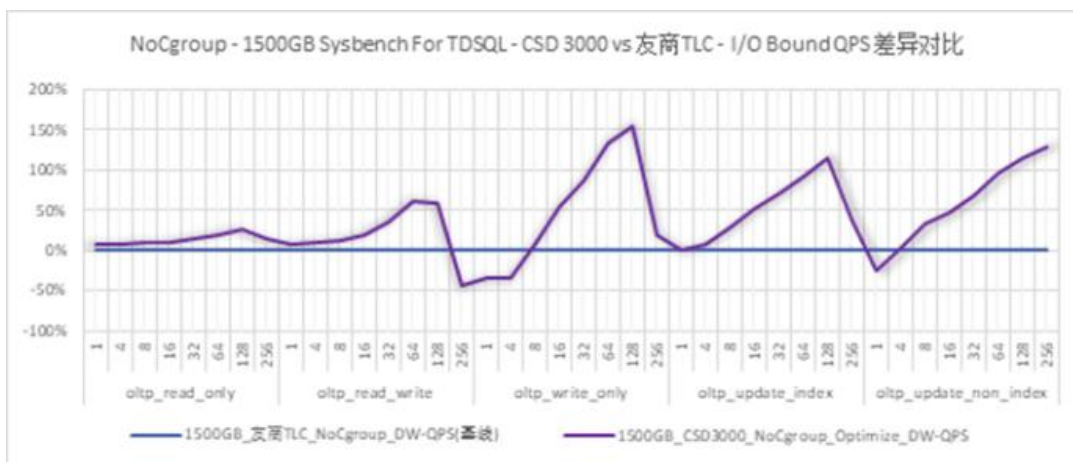


图 11： CSD 与普通 NVMe SSD 的 Sysbench For TDSQL 性能对比

目前，CSD 产品正在与众多的行业头部企业展开广泛的探索与磨合，在可计算存储的应用场景探索与落地上也小有成就：

（1）在某头部互联网出行公司，几乎所有业务系统中的 MySQL 数据库应用场景都已大量上线；

（2）在南方某头部电商公司，也在多个业务系统的 MySQL 数据库应用场景中大量上线；

（3）与国内某第一、第二云厂商也正在展开广泛的探索、磨合，且已大规模生产部署落地，相信可计算存储在不久的将来也能够在这些企业中得到大规模的应用验证。

（四）可计算存储的应用前景

随着半导体工艺技术发展步伐的逐步放缓，传统的基于 CPU 的同构计算架构越来越难以支撑计算机系统性能的持续提升。但同时以 AI/ML 与海量数据分析为代表的各种应用对计算机系统性能的需求却在不断地加速增长。这一矛盾迫使整个计算产业界从传统的同构计算架构逐步向异构计算体系架构转型，如此自然而然地孕育出了可计算存储这一新兴领域。可计算存储产品的商业化目前仍然处于初期开拓、探索的阶段。为了迈出商业成功的第一步，可计算存储产品必须同时满足三个条件：（1）能够无缝地嵌入到现有软、硬件系统生态环境内，用户应用程序不需作任何改动；（2）所提供的计算服务必须具有足够的通用性、并能带来显著的应用价值；（3）必须提供与通用存储器（如 NVMe SSD）相当的数据存储功能和性能。所以在可预见的将来，可计算存储的发展会主要体现于将透明数据压缩、透明数据加密功能集成到高性能存储控制芯片内，以保证与现有软、硬件系统生态环境之间的无缝对接，并与 NVMe 标准完全兼容。同时，NVMe 委员会正在推动 NVMe 标准的扩充，以使得用户可以更好地利用可计算存储器内部的透明压缩能力、达到降低数据存储成本的目的。目前，ScaleFlux 公司已经成功地将以透明数据压缩、加密为核心的可计算存储产品投放市场，并且多家企业正在研发类似产品。

毫无疑问，以透明数据压缩、加密为核心的可计算存储产品将会成功地实现可计算存储商业化从零到一的突破。从长期发展的角度来看，可计算存储产品所提供的计算服务远远不止透明数据压缩、加密。随着软、硬件系统生态环境逐步提升对可计算存储产品的包容性和适应性，系统会将更多的与海量数据高度耦合的计算任务（如查询、过滤、数据预处理等）直接下推至可计算存储器内，达到进一步提高整体系统计算能力和效率的目的。同时，新兴的 Chiplet 和 3D 异构集成技术也将会带来一个崭新的可计算存储器主控芯片设计空间，更进一步地提高可计算存储的价值。虽然任何人都无法准确预测可计算存储产品在接下来的十年、二十年的发展路线和轨迹，有一点是毋庸置疑的，那

就是可计算存储在未来的异构计算体系中必将占据非常重要的地位。也许在数十年之后，所有的数据存储设备都将提供多或少的计算功能服务，届时“可计算存储”这一词汇本身也就完成了其历史使命。

3、一体机

随着数字化转型的不断深入，业务需求也越来越多样化，数据量和数据类型呈现爆炸式增长。传统的关系型数据库已经无法满足需求，越来越多的企业通过采用不同类型的数据库，来满足不同应用场景的需求。这预示着多元数据库时代的来临。

一、多元数据库时代对基础设施的挑战

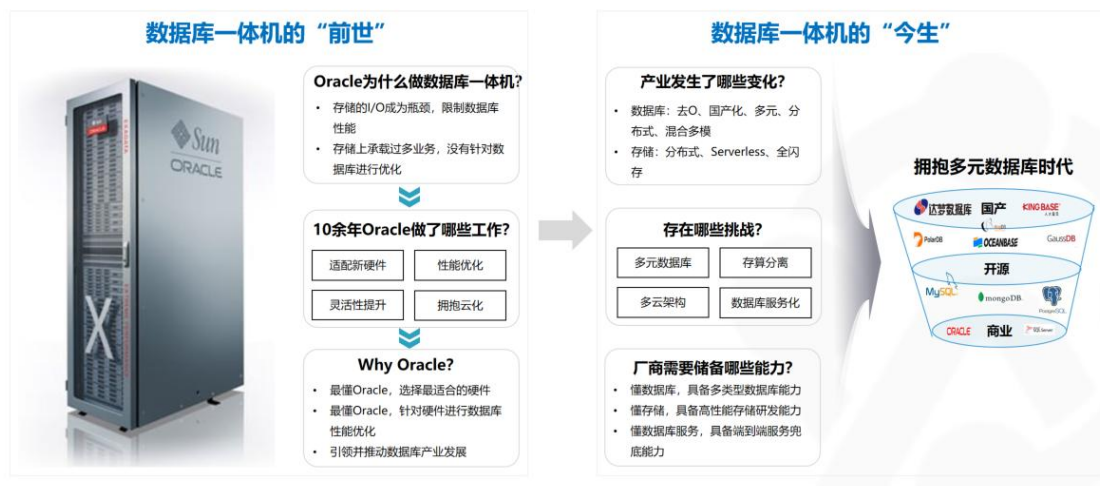
随着多元数据库技术的成熟，数据库从单一架构支持多类应用演变为多类架构支持多类应用，这些架构并非替代关系，而是相互共存、共同发展的关系，数据库的多样性成为必然，承载数据库运行的基础设施也从传统三层 IT 架构演进到各类云架构，这对基础设施也带来了众多挑战：

- **数据库类型的多样性：**要求基础设施支持多种不同类型的数据库，包括商业数据库 Oracle/DB2/SQL Server、开源数据库 MySQL/PG、国产数据库达梦/人大金仓/OpenGauss/GaussDB 等，不同类型的数据库具有不同的特性和部署方式。
- **数据库应用场景的多样性：**多元数据库需要满足不同应用场景的需求，包括 OLTP、OLAP、流数据处理等。不同应用场景对数据库的性能、可靠性等要求不同。
- **数据库管理的复杂性：**多元数据库的部署和管理更加复杂，需要采用自动化工具来提高管理效率，需要加强数据库安全管理。

二、数据库一体机的定义及发展趋势

数据库一体机作为数据库重要的基础设施之一，它是指将数据库软件和硬件服务器集成在一起，提供高性能、高可用、简单易用的解决方案。

- **数据库一体机的“前世”：**主要是 Oracle Exadata，Oracle 推出 Oracle Exadata 主要解决存储 I/O 瓶颈，对 Oracle 数据库进行性能优化，适配最新的硬件，提供灵活、弹性、云化的高性能端到端解决方案。
- **数据库一体机的“今生”：**需要主动拥抱多元数据库时代，面向数据库多元化、多云架构、国产化、数据库服务化挑战，提供全新的数据库一体机产品解决方案。

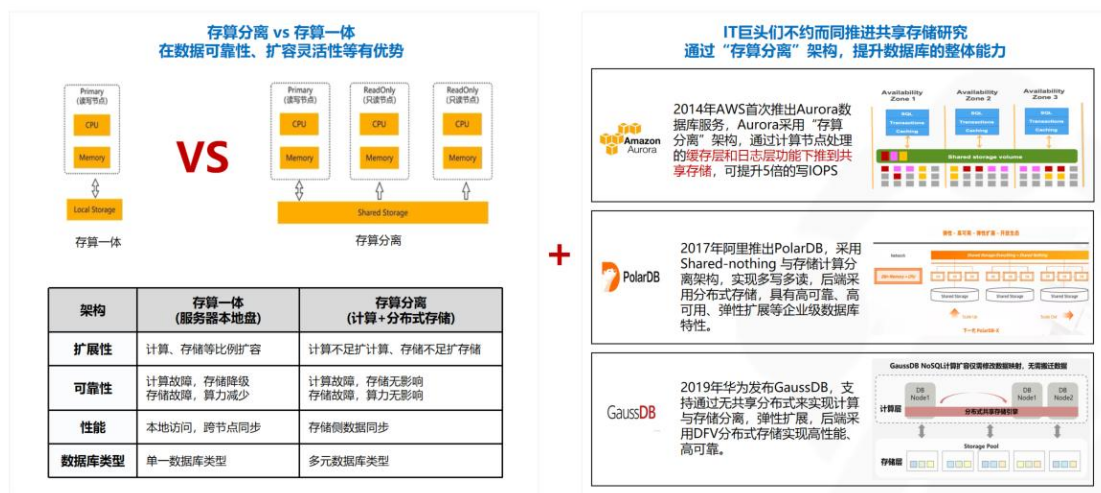


随着多元数据库时代的来临, 数据库一体机的发展, 主要呈现以下几个趋势:

- **趋势一、企业迈向多元混合数据库时代, 单一数据库一体机走向多元数据库一体机:** 数据存储结构越来越灵活多样, 新兴业务需求催生数据库及应用系统的存在形式愈发丰富。信创政策催化下, 国产数据库的发展百花齐放, 加速企业迈入多元混合数据库时代。



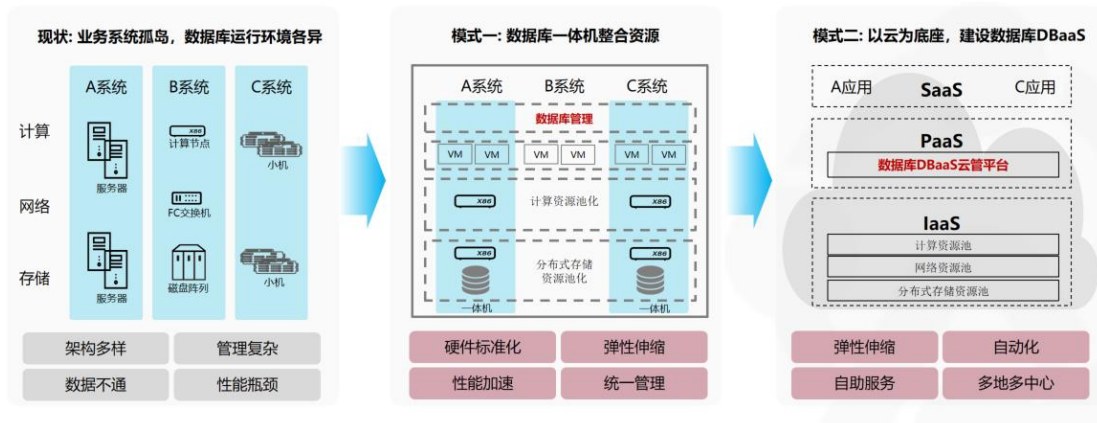
- **趋势二、数据库厂商纷纷走向“存算分离”, 采用共享分布式存储成为数据库一体机必备:** 数据库存算分离, 逐渐成为数据库的标配。存算分离后, 数据库在可靠性、主从强一致性、资源按需扩容等优势明显, 随着企业数字化转型的不断深入, “存算分离”架构正在成为企业核心数据库共同的选择, 也是最适合当前时代发展需求的一种架构。



- **趋势三、数据库厂商通过算子下推，构建独一无二的数据库优势：**随着 RDMA、FPGA 等技术的发展，存储与服务器间的网络越来越快，各大数据库厂商在考虑提升数据库整体性能时，其中一个技术方向就是将频繁需要从存储读取数据的操作，下推到存储引擎进行就近计算，这就是算子下推。但因为算子下推会涉及数据库读写逻辑的变化，目前主要还是以数据库厂商驱动为主。典型的数据库厂商包括：Oracle、PolarDB、GaussDB、腾讯 TDSQL、TiDB、AWS Aurora 等。

产品		Oracle Exadata	PolarDB	TiDB	Yugabyte DB	Cockroach DB	AWS Aurora
算子下推	Filter	√	√	√	√	√	√
	Limit	√	√	√	√	√	√
	Grouping	√	√	√	√	×	√
	Join	√	√	×	√	×	√
	Index	√	√	×	×	×	√
井置下推		√	√	×	√	×	√
日志下推		√	×	×	×	×	√
功能下推	压缩	√	×	×	×	×	×
	备份	√	√	√	×	×	×
	加解密	√	×	×	×	×	×

- **趋势四、数据库一体机具备多类型数据库安装部署、监控告警、性能分析、运维能力，实现多元数据库全生命周期管理。**传统集成架构数据库安装部署复杂、运维难，企业逐渐倾向于能提供数据库服务化的方式。目前主流的数据库服务化有 2 种部署方式
 模式一数据库一体机：构建数据库资源池，实现数据库自动安装部署、统一运维管理。
 模式二数据库 DBaaS：以云为底座，数据库云管+计算资源池+分布式存储资源池，实现云 RDS 能力。



- **趋势五、硬件层面，全球 SSD 发货维持高速增长，全闪存从 SAS 大步走向 NVMe，并驱动 CXL 创新：**伴随着数字化转型，人工智能、大数据等技术的广泛采用，SATA、SAS 逐步被 NVMe 替代，NVMe 的整体性能更好，且成本在逐年下降。Computer Express link (CXL) 技术使存储单元能够超越物理限制并利用服务器体系结构，提升数据中心效能，支持高性能 HPC 和人工智能 AI 工作负载。



SATA vs SAS vs NVMe vs CXL

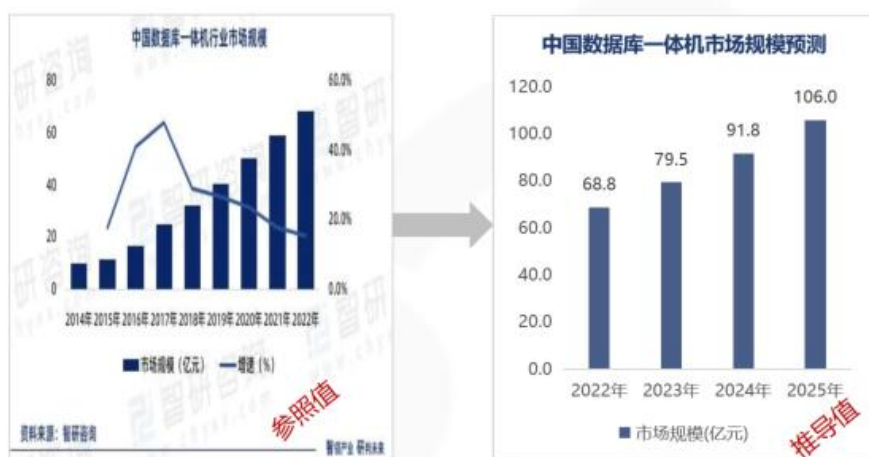
硬盘类型	SATA	SAS	NVMe	CXL
应用场景	PC、OA	虚拟化	数据库、大数据	AI、HPC
数据传输速度	600MB/s	12Gbps	32GB/s	32 GT/s
IOPS	10万	40万	100万	2000万
延迟	ms	us-ms	us	us
容量	4-14TB	300GB-1.2TB	19.2-7.68TB	128-512GB
每Gb价格	低	一般	较高	高
Multi-channel	依赖第三方控制器	依赖第三方控制器	支持	支持
直连内存	不支持	不支持	支持	内存扩展

三、数据库一体机市场分析

随着中国数据库市场的扩大，以及数字化转型和数据驱动业务的普及，传统的 IT 技术架构逐步地向云化、国产化演变，企业对于高性能的数据处理和存储解决方案的需求增加，用户对数据中心基础设施的高安全、高性能、高稳定等需求提升到了一个新的高度，这也推动了数据库一体机市场的蓬勃发展。

据智研咨询测算，2022 年中国数据库一体机市场规模约为 68.8 亿元，年增长率为 15.5%。按推导，如果按照年增长率 15.5%进行测算，预计到 2025 年中国数据库一体机市场规模约为 106 亿元。

中国数据库一体机市场规模及预测



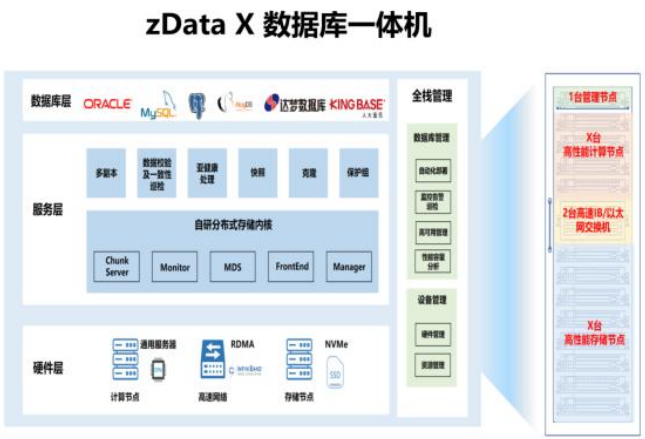
- 目前市面上，数据库一体机产品繁多，可以按照应用场景、部署模式、厂商、数据库类型等进行分类。其中在应用场景上，OLTP 仍然是主流，占比 80%以上。



- 数据库一体机的厂商玩家众多，基本分为四类：数据库生态厂商、数据库厂商、服务器厂商和存储厂商，这些厂商各自发挥各自的优势。



- **数据库生态厂商：**云和恩墨多元异构数据库一体机 zData X，兼容 Oracle、MySQL、SQLServer、OpenGauss、人大金仓、达梦等数据库。支持全栈信创硬件，包括鲲鹏/海光芯片、麒麟操作系统等，以及提供原厂维保服务和数据库专业服务。



- **数据库厂商：**达梦 DAMENG PAI 数据库一体机，内置 DM8，为用户提供“高性能、高可用、

易维护、可扩展”的数据库服务，以及提供原厂维保服务、提供达梦数据库专业服务。



- **服务器厂商：**超聚变 FusionOne 数据库一体机，基于自主研发的 FusionServer 服务器，与 Oracle 完美适配，使其具备高性能、高可用、易运维的特点，覆盖了性能调优、便捷部署、系统管理等能力，为用户提供企业级数据库解决方案。



- **存储厂商：**华为 FusionCube 超融合一体机，预集成分布式存储引擎、虚拟化和云管理软件，资源可按需调配、线性扩展，支持 SAPHANA，Oracle 等。



四、企业客户如何评估选择数据库一体机？

数据库一体机的选型, 需要从数据库一体机全生命周期视角对厂商品牌、产品能力、服务角度, 可以从怎么选、怎么替、怎么管角度综合评估：

- **怎么选：**①选具备自主可控能力，且在数据库乃至数据库基础设施都有研发积累沉淀的品牌。②针对产品选择高性能、高可靠、多数据库兼容且支持信创部署的产品。③选提供产品原厂维保服务，且能提供数据库专业咨询运维服务的厂商。
- **怎么替：**①避免采购绑定，可软硬解耦部署。②产品易安装部署、厂商提供数据迁移能力。③可提供本地化服务。
- **怎么管：**①需具备多类型数据库安装部署、管理能力。②具备硬件设备资源管理能力。③除了产品管理能力，亦可提供技术支持及数据库专业服务能力。

数据库一体机 评估维度		怎么选？	怎么替？	怎么管？
品牌	- 是否自主可控？ - 是否在数据库领域有影响力？ - 是否在基础设施领域有影响力？	是否采购绑定？		
产品	- 性能：是否满足数据库高性能？ - 可靠：是否提供多重可靠保障？ - 扩展性：是否按需扩展？ - 兼容性：是否兼容多类型数据库？ - 架构演进：是否兼容现在&未来演进？ - 国产化部署：是否支持国产架构？	- 安装部署：是否易安装易部署？ - 数据迁移：是否具备迁移能力？ - 投资保护：是否能保护已有投资？	- 是否提供数据库管理？ - 是否提供硬件资源管理？	
服务	- 是否原厂提供产品维保服务？ - 是否能提供数据库专业服务？	是否有本地化交付服务？	- 是否提供技术支持服务？ - 是否提供数据库专业服务？	

另外，企业在评估选择数据库一体机时，还可以参考以下几点：

- 企业可以咨询专业人士，获取更专业的建议。
- 企业应该多家对比，选择最适合自身需求的数据库一体机，确保未来演进到国产数据库。
- 企业需要结合自身采购模式，选择解耦采购、避免绑定。

4、智能运维

随着人工智能、云计算、大数据等新技术的快速发展，业务需求越来越多样化，以及伴随国产数据库的不断完善和成熟。越来越多的企业通过采用不同类型的数据库，来满足不同应用场景的需求。数据库从单一架构支持多类应用演变为多类架构支持多类应用，这些架构并非替代关系，而是相互共存、共同发展的关系，数据库的多样性成为必然，这预示多元数据库运维管理时代全面来临。

一、数据库运维挑战与趋势

数据库运维管理越来越复杂，面临着诸多挑战。这给日常运维数据库的 DBA 们带来了数据库运维管理挑战，主要包括：

- **数据库类型多、数量大，管理复杂：**企业采用多种不同类型的数据库来满足不同的业务需求，不同类型的数据库有不同的运维方式和工具，海量数据的管理和维护非常复杂和耗时。
- **数据库专业知识要求高，DBA 们无法短时间掌握多类型数据库知识：**不同类型的数据库设计实现及运维方式有所不同，DBA 疲于日常运维，很难同时在多种数据库上专业精通。
- **业务变化快、需求多样，无法快速满足业务对数据库的变更要求：**企业的数据库使用场景也越来越多样化，包括业务应用、分析应用、IoT 应用等，不同使用场景的数据库具有不同的需求，如何针对跨类型应用场景做出快速响应带来了挑战。
- **数据库运维成本高昂：**数据库运维是一项非常耗费人力和物力的工作。随着数据库规模的不断增长和安全威胁的日益严峻，数据库运维成本也变得越来越高昂。

为了应对以上挑战，企业可以采取先建数据库管理平台的思路进行应对，采用统一的数据库管理平台，通过自动化工具和智能化技术，帮助企业统一管理多种不同类型的数据库，实现数据库运维的自动化，提高运维效率。

目前市面上数据库运维管理软件分门别类，但总体呈现以下趋势：

- **多源化**：随着数据库的多元化，数据库管理需要支持更多的数据源，包括关系型数据库、NoSQL 数据库、时序数据库、云数据库等。
- **自动化**：自动化将成为数据库运维管理的重要趋势。自动化可以帮助减少人工干预，提高运维效率和可靠性。例如，自动化部署可以帮助快速部署数据库，自动化监控可以帮助及时发现数据库问题。
- **智能化**：人工智能技术正在被广泛应用于数据库运维管理领域。人工智能驱动的数据库运维工具可以帮助运维人员更快、更准确地识别和解决问题。人工智能还可以帮助运维人员预测数据库的性能瓶颈，并进行提前干预。例如，机器学习可以用于预测数据库故障，人工智能可以用于优化数据库性能。
- **多云化**：随着云计算技术的快速发展，企业越来越多地采用多云策略，数据库多云管理变得越来越重要。多云数据库管理平台可以帮助企业在多个云平台上管理和操作数据库，提供统一的管理界面、自动化运维工具和安全合规功能，帮助企业提高跨云数据库的可用性、性能、成本和敏捷性。
- **协同化**：传统的数据库运维管理模式往往是各部门各自为政，缺乏协同。这不仅会导致运维效率低下，还会增加数据库故障的风险。协同将成为数据库运维管理的重要趋势。通过跨团队合作、统一的运维流程、共享工具和资源、知识共享等措施，提高运维数据库的高可用性、性能和安全性。例如，DevOps 模式下强调开发团队和运维团队的紧密合作。

二、数据库运维管理新技术展望

（一）展望一、AIGC 技术逐步应用到数据库运维管理

随着 AIGC 技术的不断发展，AIGC 在数据库运维管理中具有广阔的应用前景，可以帮助数据库运维人员提高工作效率、降低成本并提高服务质量。AIGC 在数据库运维管理中的具体应用包括：



- **智能故障诊断：**AIGC 可以帮助数据库运维人员快速准确地诊断数据库故障。AIGC 可以分析数据库日志、性能指标和其他数据，并自动识别故障的根本原因。这可以帮助数据库运维人员更快地解决问题，减少数据库宕机时间。
- **自动性能优化：**AIGC 可以帮助数据库运维人员自动优化数据库性能。AIGC 可以分析数据库负载、资源利用率和其他数据，并自动调整数据库配置参数和索引策略。这可以帮助数据库以最佳性能运行，并满足业务需求。
- **智能问答系统：**AIGC 可以帮助搭建智能问答系统，帮助企业用户快速、准确地获取数据库运维相关的知识和信息，提高数据库运维效率和效能，包括数据库故障诊断和处理、数据库性能优化、数据库安全防护等。
- **智能容量规划：**AIGC 可以帮助数据库运维人员进行智能容量规划。AIGC 可以分析数据库负载、增长趋势和其他数据，并自动预测数据库未来的容量需求。这可以帮助数据库运维人员提前规划数据库的扩容，避免数据库出现性能瓶颈。

（二）展望二、数据库运维服务经验与数据库运维管理软件的结合会更加紧密

- 数据库运维服务经验和数据库运维管理软件可以形成能力循环，相互促进，共同提高数据库运维管理水平，具体体现在以下几个方面：



- **数据库运维服务经验反哺数据库运维管理软件：**数据库运维服务人员在实际工作中积累了丰富的经验，这些经验可以帮助数据库运维管理软件开发商改进软件功能，提升软件的运维能力。例如，数据库运维服务人员可以通过分析数据库故障案例，帮助软件开发商开发更准确的故障预测模型。
- **数据库运维管理软件反哺数据库运维服务：**数据库运维管理软件可以为数据库运维人员提供自动化、智能化的运维工具，帮助数据库运维人员提高运维效率和准确性。例

如，数据库运维管理软件可以自动收集数据库运行数据，帮助数据库运维人员发现数据库运行问题。

- **数据库运维服务与数据库运维管理软件共同提升数据库运维能力：**数据库运维服务人员和数据库运维管理软件可以相互配合，共同提升数据库运维能力。例如，数据库运维服务人员可以利用数据库运维管理软件收集的数据，进行数据库故障分析和性能优化。

（三）展望三、跨云数据库运维管理将成为主流

多云将成为未来企业主流架构，跨云数据库运维管理将成为企业实现多云战略的必要条件。通过采用合理的策略和工具，企业可以有效地进行跨云数据库运维管理，确保数据库的稳定运行和业务的持续发展，具体包括以下几点：



- **构建跨云数据库管理平台：**企业采用功能全面、兼容性强的跨云数据库管理平台，不仅仅是云数据库，同时兼顾私有云环境传统数据库，提供统一的管理界面、自动化运维工具和安全合规功能。
- **建立统一的跨云数据库管理规范 and 流程：**统一的跨云数据库管理规范 and 流程可以帮助企业解决管理挑战。企业应根据自身需求制定统一的跨云数据库管理规范 and 流程。
- **建立跨云数据库运维团队：**负责跨云数据库的日常运维和故障处理。跨云数据库运维团队需要掌握不同云平台的运维工具和方法，并能够熟练使用跨云数据库管理平台。
- **加强数据安全防护措施：**企业应采用安全措施，保障数据在传输过程中的安全。例如，企业可以采用数据加密、身份认证等措施，保障数据安全。

三、数据库运维管理软件现状与不足

随着数据库使用量的不断增加，数据库运维管理软件市场规模也在不断增长。根据中国信通院《2023 年中国数据库运维管理软件市场研究报告》，中国数据库运维管理软件市场规模将从 2022 年的 30.5 亿元增长到 2027 年的 73.2 亿元，年复合增长率为 18.9%。

虽然伴随数据库运维管理市场的蓬勃发展，数据库运维管理软件的功能日益丰富，数据库运维管理软件的功能日益丰富，包括数据库备份、恢复、监控、性能优化、安全防护、数据管理等。

但目前中国市场整体数据库运维管理软件也存在以下不足：

- **兼容能力弱**：许多数据库运维管理软件只支持仅有的几款数据库，或只覆盖云上开源数据库，或只覆盖云下传统数据库，无法满足企业用户面向未来国产化转型，即覆盖当下，又满足未来演进的运维管理诉求。
- **专业性不强**：许多数据库运维管理软件仅仅充当监控工具的角色，简单地做数据采集和展示，没有融入专业的数据库运维管理方法论或最佳实践，实践中具体问题还得依靠 DBA 进行人工查看和分析。
- **人工智能技术应用低**：数据库运维管理软件的智能化水平不高，AIGC 与数据库的结合大多处于研究阶段，难以对数据进行深度分析，做出准确的判断。
- **难以满足企业个性化需求**：许多数据库运维管理软件无法满足企业的个性化需求，需要企业进行大量的二次开发，无论客户还是厂商耗时耗力，效果却不凸显

四、企业客户如何评估选择数据库运维管理软件？

站在企业 DBA 的角度，选择数据库管理工具，可以侧重评估以下几方面：

- **类型**：企业需要考虑自身使用的多种数据库类型，确保数据库运维管理软件能够支持这些数据库类型。
- **功能**：数据库运维管理软件应该具备基本的功能，如数据库监控、故障诊断、性能优化、安全管理等。
- **性能**：数据库运维管理软件需要能够满足企业的性能需求。
- **易用性**：数据库运维管理软件应该具有良好的易用性，能够方便运维人员使用。
- **价格**：需要根据企业的实际情况进行价格考虑。
- **厂商**：数据库运维管理本质是数据库运维服务经验产品化，需要考虑厂商是否具备丰富的数据库运维经验积累。
- **售后服务**：企业需要选择能够提供优质售后服务的厂商的软件。

目前市面上有很多类型的数据库管理工具，都各有千秋。常见的数据库管理软件包括：

- **云和恩墨 zCloud 数据库云管平台**：面向多元混合数据库环境，提供跨多云架构、跨多类数据库的一站式智能数据库管理平台。持续汇聚专家知识和经验，融合行业标准

和最佳实践，通过多元数据库统一纳管，实现标准化、自动化、智能化的数据库全生命周期管理。



- **沃趣 QFusion 容器化数据库发放平台：**采用 Kubernetes 等云原生技术，提供云上容器数据库一站式的数据库管理服务。
- **阿里云 DMS 数据管理平台：**提供云上全域数据资产管理、数据库设计开发、数据集成、数据开发及数据消费等功能。

另外，企业在评估选择数据库运维管理软件，还可以参考以下几点：

- 企业可以参考《中国信通院数据库管理平台专项评测》进行评估，评测采用了国际标准的评测方法，面向国内外数据库厂商和服务商的数据库管理平台进行功能、性能、安全、可靠性等方面的评测，评测结果具有权威性和公正性，旨在为企业用户提供数据库管理平台选型参考。
- 企业可以咨询专业人士，获取更专业的建议。
- 企业应该多家对比，选择最适合自身需求的数据库运维管理软件。
- 企业应该先试用数据库运维管理软件，确保其能够满足自身需求。

5、云管平台

由于数据库云化、多元化、信创化等趋势的影响，数据库管理范式已经发生了巨大的变化。

一、数据库多元化

数字化场景在不断丰富，企业经常需要选择多种数据库来满足不同业务场景的需求，如大中型企业通常会使用多达十多种类型的数据库，包括各种关系型数据库，以及图、文档、时序、宽表、搜索引擎等非关系型数据库。根据 DB-Engines 的最新数据，目前全球共有近 400 种开源和商业数据库。而如何对这些多元化的数据库进行统一管理，就成了困扰企业的一个难题。

二、数据库云化

传统的在本地数据中心部署和使用数据库的模式，其成本高且非常低效，导致越来越多的企业逐渐开始选择更加敏捷、弹性，以及成本更低的云端部署的数据库。Gartner 预测，在 2025 年将有 75% 的数据库运行在云端，云端数据库的 DBaaS 使用模式已经极大地影响了线下数据库的使用范式，DBaaS 背后是一种服务化的思维，屏蔽了各种数据库的复杂度，以及数据库底座的差异性和复杂度，极大地简化了数据库的管理，让业务人员有更多的精力专注于业务创新及架构设计等高级工作。

三、数据库信创化

国内信创政策也导致了数据库多元化和 IT 架构复杂的问题进一步加剧。在国内市场，金融、政务、能源等行业企业出于满足信创政策的需求，在很多业务系统引入多款国产数据库，如 GaussDB、Oceanbase 等，这进一步加深了企业内部数据库多元化的局面。而在底层资源方面，国产需求用户会使用基于鲲鹏、海光、飞腾等多种芯片架构的服务器，以及使用麒麟、统信等多种国产操作系统。面对国产硬件+系统+数据库的乘法效应，让企业的基础软硬件与 IT 业务环境变得更加复杂，加大管理难度。

四、企业级能力要求高

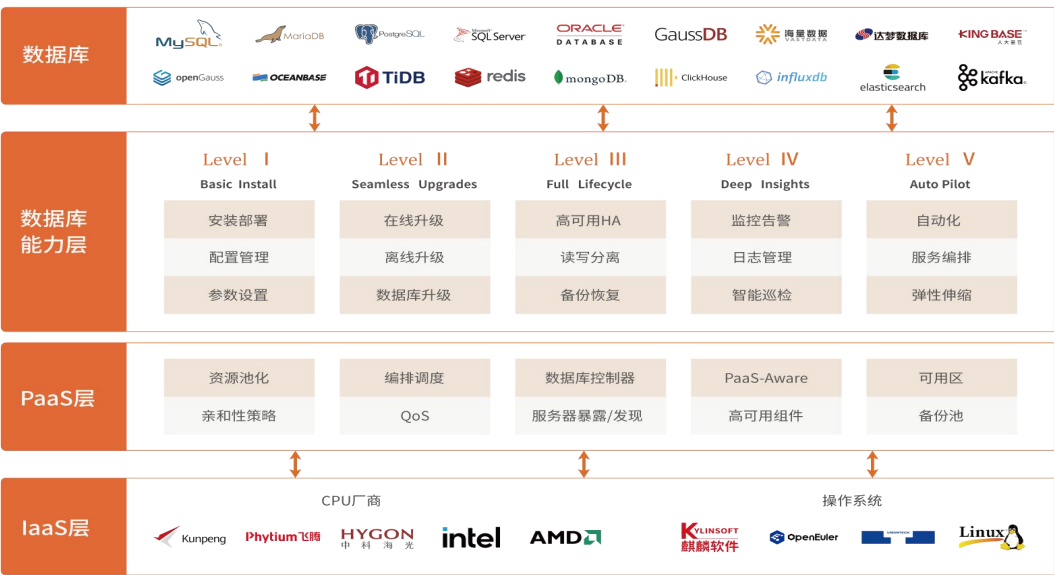
数据库的性能、可靠性等企业级能力，往往关乎国计民生，大量的行业业务特点，都对数据库的 RTO、RPO 等指标有着极高的要求，因此企业级的 dbPaaS 需要能够提供多 Region 容灾、跨 AZ 双活、高密度数据整合等能力。

新一代的数据库统一云管 dbPaaS 平台，能够很好地解决以上多个问题。例如沃趣科技的 QFusion 数据库私有云平台，基于新一代云原生技术，能够为 18+ 种以上全球主流数据库提供全生命周期的管理能力，其中也已经支持国产 GaussDB、Oceanbase 等七款国产数据库。QFusion 完全基于 DBaaS 服务化、云化的设计理念，彻底释放业务人员、运维人员使用数据库、管理数据库的成本，通过产品服务目录，可以以极低的学习成本，完成绝大多数的管理运维工作。QFusion 打造了强大 QPaaS 层，向上 QPaaS

能够对数十种数据库提供编排调度、资源隔离、亲和/反亲和、服务暴露等能力，向下能够屏蔽底层资源的差异性和复杂度，能够对资源进行统一抽象，做到对资源池化，PaaS-Aware，相当于专为数据库打造的分布式操作系统。此外，QFusion 提供了多 Region 容灾、跨 AZ 双活、高密度数据整合等企业级能力，可以为金融级场景的用户提供企业级的数据库 PaaS 平台，有效降低数据库运维与管理成本。

QFusion数据库私有云RDS

支持18款全球及国产主流数据库，实现数据库的全生命周期管理能力



6、安全管控平台

一、发展背景及概述

在混合云与国产化风潮的推动下，数据库行业迎来了蓬勃发展的时机。随着各类企业数字化转型的加速进行，数据库量呈现出爆发性增长的趋势。这种增长不仅仅体现在数据规模的扩大，还表现在多样性和复杂性上。新兴技术如 NewSQL、分布式系统、HTAP（混合事务/分析处理）、Serverless 架构、湖仓一体化、内存技术以及云原生模型等不断涌现，构筑了一个庞大而复杂的数据库生态系统。

在这个多元化的数据库环境中，企业面临着前所未有的挑战。其中之一是如何高效地管理和操作异构数据库。企业可能同时使用关系型数据库、NoSQL 数据库、内存数据库等多种类型，需要统一的管理和操作手段来确保数据的一致性和高效性。

同时，在数据爆发背景下，数据库作为数据存储与处理的核心，时刻都在应对海量数据。安全隐患的存在可能直接影响到核心业务的正常运行，甚至威胁社会稳定。据近

年来数据库安全事件的统计显示，数据库安全面临的重大威胁来自内部。恶意内部人员和人为错误是对数据库安全最大的挑战，而非外部入侵者。因此，对内部安全的关注至关重要，需要对数据库的访问人员和操作行为进行有效管控，避免因内部原因导致的数据泄漏、数据破坏和数据库不可用等问题。

在此背景下，实现高效且安全的一体化操作管控成为企业的重要任务。企业不仅需要应对不同行业的背景和业务问题，还需逐一解决数据库管理与安全管控的复杂挑战。本主题旨在通过分享企业在一体化数据库管控建设方面的实践经验，为其他企业在应对多样化数据库和不断升级的数据安全标准时提供有益见解和实用经验。通过深入探讨解决方案、技术实现以及成功案例，促使企业更好地理解如何在这个瞬息万变的数据库生态中取得平衡，实现业务的可持续发展。

二、一体化数据库管控平台解决的问题

在上述背景下，企业面临着较多数据库管理的效率问题和安全问题：

（一）效率层面

- **性能瓶颈：** 数据库性能问题可能导致查询响应时间增加，从而影响业务流程。优化数据库性能需要监控和调整数据库配置、查询优化以及索引设计。
- **备份与恢复效率：** 数据备份和恢复的过程可能耗费大量时间和资源，特别是在处理大规模数据时。高效的备份策略和恢复机制对于业务连续性至关重要。
- **版本升级与迁移难题：** 数据库系统的版本升级和平台迁移可能面临复杂的问题，包括数据迁移、应用程序兼容性等，这可能导致业务中断和系统不稳定。
- **多数据库平台管理复杂性：** 企业通常使用多个数据库平台，管理和协调这些异构数据库系统可能变得复杂，需要额外的资源和工具。

（二）安全层面

- **内部威胁：** 恶意内部人员或不当操作可能导致数据泄露、篡改或删除。
- **外部攻击：** 数据库可能成为外部黑客攻击的目标，尤其是包含敏感信息的数据库。
- **弱密码和访问控制不当：** 弱密码和不合理的访问控制策略可能导致未经授权的访问。
- **缺乏加密：** 数据库中的敏感数据未经适当加密可能面临泄露的风险。
- **审计和监控不足：** 缺乏对数据库操作的审计和监控可能导致对安全事件的迅速察觉困难。

一体化数据库管控平台通过集成和统一管理多个数据库系统，能够实现数据管理的集中化和自动化。例如，通过一键式的数据库部署和配置功能，管理员可以快速、轻松地部署和配置多个数据库实例，大幅缩短了传统手动部署的时间，提高部署效率。此外，这类平台能够提供统一的监控和管理界面，管理员可以通过该界面实时监控各个数据库的运行状态和性能指标，快速定位和解决问题。

在安全层面，一体化数据库管控平台基本配备多层次的安全防护机制，以此来保障数据的安全性和完整性。例如，平台支持对数据库的访问权限进行精细化控制，管理员可以根据用户角色和权限需求，灵活设置数据库的访问权限，防止未授权的用户访问和操作数据库。同时，平台还提供了安全审计功能，记录和分析用户的操作行为和数据库的变更情况以及时发现和应对安全威胁，保障了数据库的安全性。此外，一体化数据库管控平台能够支持数据加密和脱敏功能，对敏感数据进行加密和脱敏处理，有效防止数据泄露和信息泄露风险，对企业资源管控提供了很大帮助。

三、一体化数据库管控平台的实现原理

一体化数据库管控平台的实现原理基于现代化数据库管理技术和信息安全理论，通过软件系统和硬件设备相结合的方式，实现了对数据库的统一管理、安全保护和性能优化，为企业提供了全方位的数据库管理和安全保障。

其实现原理主要包括以下几个方面：

（一）数据管理与监控技术

平台通过数据库管理系统（DBMS）和监控系统实现对数据库的管理和监控。数据库管理系统用于管理数据库的结构和内容，包括表空间管理、索引管理、事务管理等；监控系统用于实时监测数据库的运行状态和性能指标，包括 CPU 利用率、内存利用率、磁盘 IO 等。

（二）安全防护与权限控制技术

平台通过安全防护系统和权限控制系统实现对数据库的安全防护和权限控制。安全防护系统包括数据加密、访问控制、身份认证等技术，用于保护数据库中的敏感数据；权限控制系统包括用户权限管理、角色权限分配等技术，用于控制用户对数据库的访问权限。

（三）性能优化与自动化运维技术

平台通过性能优化系统和自动化运维系统实现对数据库的性能优化和自动化运维。性能优化系统包括查询优化、索引优化、资源调度等技术，用于提升数据库的查询效率和响应速度；自动化运维系统包括任务调度、故障处理、日志记录等技术，用于实现数据库的自动化运维和管理。

（四）数据备份与灾备恢复技术

平台通过数据备份系统和灾备恢复系统实现对数据库的备份和灾备恢复。数据备份系统包括全量备份、增量备份、定时备份等技术，用于保障数据库数据的安全性和完整性；灾备恢复系统包括故障恢复、灾难恢复、数据恢复等技术，用于实现对数据库的灾备和恢复。

（五）信息安全与合规监管技术

平台通过信息安全系统和合规监管系统实现对数据库的信息安全和合规监管。信息安全系统包括数据加密、安全审计、风险评估等技术，用于保障数据库的信息安全；合规监管系统包括合规审计、报告生成、监管报告等技术，用于实现对数据库的合规性监管和报告生成。

四、一体化数据库管控平台的应用案例

一体化数据库的好处不言而喻，接下来我们来看一个金融行业 TOP 企业案例。该大型金融机构拥有庞大的数据资产，使用的数据源繁多，涉及客户信息、交易记录、风险管理等敏感数据，业务场景十分复杂，对此需要进行全面而严格的管理和保护。

（一）痛点与需求

- 1. 数据分散管理：** 该金融机构拥有多个业务系统和数据库，数据分散在不同的平台和环境，导致管理困难和效率低下。
- 2. 数据安全保护：** 由于金融数据的敏感性，需要对数据进行严格的安全保护，防止数据泄露和非法访问。
- 3. 合规性要求：** 金融行业受到监管机构的严格监管，需要满足合规性要求，对数据库管理和数据处理流程进行规范和监控。

（二）解决方案

该金融机构选择了一体化数据库管控平台作为解决方案，以应对上述痛点和需求。

- 1. 统一管理平台：**通过一体化数据库管控平台，在不同数据库之间建立连接和桥梁，实现了对多个数据库系统的统一管理，包括 Oracle、SQLServer、MySQL 等，提供一致性的管理和监控接口，统一的管理界面和操作流程大大提高了管理效率。
- 2. 强大功能：**平台通过采用先进的技术，如 Agent 技术、轻量级代理、API 调用等，实现了对各类数据库的实时监控、性能优化、备份恢复、安全管理等功能。
- 3. 安全防护机制：**一体化数据库管控平台提供了多层次的安全防护机制，包括访问控制、数据加密、安全审计等功能，有效保护了金融机构的数据安全。
- 4. 合规性监管：**平台提供了丰富的监控和审计功能，可以对数据库的访问和操作进行全面监控和记录，确保数据库管理和数据处理流程符合监管机构的合规要求。

（三）成效剖析

- 1. 提升效率带来的成本节约：**通过提升数据管理效率，该金融机构节省了大量的人力资源和时间成本。数据库管理员的工作效率得到了显著提升，数据库管理工作变得更加轻松高效。
- 2. 数据安全保护的效果显著：**平台提供的安全防护机制有效保护了金融机构的数据安全。通过加密、访问控制等手段，成功防止了数据泄露和非法访问等安全风险的发生，为金融机构的数据安全提供了坚实保障。
- 3. 合规性要求的达标保障：**平台提供的合规性监管功能帮助金融机构规范了数据库管理和数据处理流程，确保了数据处理符合监管机构的合规要求。金融机构能够及时发现和纠正不符合规定的行为，降低了合规性风险。

7、SQL 审核

一、概述

在过去的几年中，随着国产数据库的迅猛发展，对数据库行业带来的变化让国内企业对数据管理的需求日渐增长。SQL 审核工具已经从基本的监控工具演变为提供综合数据治理、风险管理和合规性支持的复杂解决方案。企业对这些工具的依赖日益增加，这推动了相关产品的创新和发展。

二、趋势

（一）国产化浪潮：数据库管控新考验

国产数据库的兴盛已然成为行业中的一个主流趋势。随着数据库产品日益增多和快速迭代，目前尚未建立起一套统一的操作规范和成熟的最佳实践。这种现状给技术产业的各个层面带来了显著的挑战。因此，企业迫切需要掌握正确使用数据库 SQL 的知识，以满足日益增长的数据库使用需求。

首先，SQL 管控工具和服务必须适应国产数据库产品的特性。这意味着审核工具需要对国产数据库的架构、性能特征、安全机制等有深入地理解和适配，以确保审核的有效性和准确性。

其次，国产数据库尚未经历足够多的实战磨炼，这要求 SQL 管控平台能够帮助企业更好地发现潜在的问题和风险。管控工具不仅要检查标准的 SQL 安全和性能问题，还要能够针对国产数据库可能存在的特殊问题进行诊断和建议。应该从原来由客户定义什么是自家企业的最佳实践规范，演进为平台帮助客户制定最佳实践。

总的来说，数据库国产化对 SQL 管控工具提出了更高的要求，这既是挑战也是机遇。对于开发商而言，这是一个需要不断创新和适应新环境的时期。对于使用者来说，在当前的关键节点，通过平台和工具的正确选用，以帮助新数据库快速、高效且低风险地使用是至关重要的。

（二）工具之上：超越工具的新篇章

首先，是安全和合规性需求的扩展。过往受限于数据保护法规变得愈加严苛，这迫使 SQL 审核工具的功能必须扩展，以确保符合这些法规。工具不再是单一的查询执行程序，它们需要能够综合性地确保企业的数据库活动全面遵循安全合规性标准。

其次，随着 DevOps 多年的实践发展，已经形成了稳定成熟的体系，得益于 DevOps 强大的支持体系，SQL 审核也迈向了新的阶段，从繁琐的行政手动作业转变为追求自动化，从单一的审核工具转变成能在事前风险防范和日常治理，事中快速定位问题和故障，事后能持续优化和改进的全套方案。

这样的变革不仅保障了数据库操作的安全性和合规性，而且还与快速迭代的要求相符，显著加快了软件的开发和部署周期。

此外，企业对 SQL 访问入口的需求正在发生转变。原本仅限于基础的数据库查询和访问，现在已经扩展到对更高级功能的迫切需要。企业正寻求通过一个统一的解决方案，无论是将现有的多种 SQL 工具集成在一起，还是创建一个独立的统一访问点，来增强数据访问和治理能力。这包括但不限于：提升安全性和合规性、强化权限管理、实现全面审计以及标准化流程等关键方面。

综上所述，从单一的工具到全面的解决方案，这一趋势反映了市场对于能够提供全方位服务的数据库管理系统的渴望。随着这些需求的增长，我们可以预见到，未来的数据库管理解决方案将会更加智能、集成化，并且能够为企业提供更多的价值。

（三）产品案例

1. 云和恩墨 SQM

SQM 是云和恩墨开发的 SQL 质量管控软件，能够自动抓取数据库开发与运行环境中的对象设计与 SQL 信息，并依据既定的审核规则分析，找出对象设计与 SQL 中的潜在问题，给出专业改进建议，规避应用性能和稳定性风险。SQM 可在应用开发、测试、上线、生产不同阶段对 SQL 进行质量管控，前置性地保障应用稳定、高效运行。

SQM 基于云和恩墨自研独有的 SQL 解析引擎，帮助企业解决多种数据库 SQL 统一解析的问题,通过对静态 SQL 和动态 SQL 的解析，包括 DDL、DML、DQL、DCL 等语句，从多个维度进行 SQL 分析，并判断 SQL 执行目的，更加智能。其核心功能如下：

（1）数据库审核

数据库审核通过连接数据库，采集 SQL 信息、数据字典信息，对 SQL 语法进行审核，并给出审核问题和建议，用户可结合执行计划和相关对象做对应的优化操作，主要应用于测试阶段的数据库审核和生产阶段的问题 SQL 审核。

（2）应用程序审核

应用程序审核对应用程序中执行的 SQL 进行分析评估。目前支持 Oracle、MySQL、DB2、PostgreSQL。通过 Java 探针实时抓取应用程序执行的 SQL 信息，根据既定规则进行分析，找出其中存在风险的 SQL。

（3）SQL 脚本审核

SQL 脚本审核将 SQL 脚本(DML 和 DDL)提交到平台进行分析评估，SQM 根据审核用户提供的 SQL 脚本，形成问题 SQL 报告。

（4）工单审核

用于管理 SQL 处理的流程。用户将 SQL 脚本提交到平台、并指定相关负责人进行审批。平台根据既定规则，对提交的 SQL 脚本进行预审核，标记其中违反的开发规范/安全规范/变更规范等。用户也可以向 DBA 申请协助处理 SQL,然后流转给审批进行复核。复核通过后，系统自动进行后续处理，如执行相关 SQL 脚本、生成数据导出文件等。

(5) Jenkins 插件的 SQL 审核

Jenkins 插件能够自动分析代码仓库中 mybatis 文件和变更脚本文件，实现 SQL 的自动审核。

(6) OpenAPI

以 RESTful 方式开放 SQL 审核能力及平台内的数据。

(7) 智能优化

分析优化通过分析 SQL 语法树、数据字典以及统计信息，给出 SQL 优化建议。

2.爱可生 SQLE

云树 SQLE 是爱可生自主研发支持多元数据库的 SQL 质量管理平台，应用于开发、测试、上线发布、生产运行阶段的 SQL 质量治理，通过“建立规范、事前控制、事后监督、标准发布”的方式，为企业提供 SQL 全生命周期质量管控能力，规避业务 SQL 不规范引起的生产事故，提高业务稳定性，也可推动企业内部开发规范快速落地。其核心功能如下：

(1) 丰富的数据源支持

支持主流商业和开源数据库，包括 MySQL、PostgreSQL、Oracle、SQL Server、DB2、TiDB、OceanBase、ActionDB、达梦、GoldenDB 等，持续增加新的数据源类型，以满足您不同的需求。

(2) 全面的审核规则

审核规则源自我们经验丰富的 DBA 运维专家团队多年的技术积累。您可以借助我们的审核规则快速达到专家级的 SQL 诊断能力，并且我们的规则库目前已经拥有规则 700+，并在不断增加中。

(3) 智能的 SQL 采集

提供多种智能扫描方式，例如慢日志、JAVA 应用等，以满足您的事前和事后 SQL 采集需求。一旦配置完成，我们的系统会自动持续采集各业务库中的 SQL，极大地减轻了您对 SQL 监督的压力，并及时发现问题 SQL。

(4) 高效的审批流转路径

提供标准化的工作流，化解了在沟通和进度追踪上的难题，从而提升了上线效率。您可以通过与飞书、钉钉等多种消息通道的对接，及时了解更新进度，减少了沟通交流的成本。

(5) 便捷的 SQL 数据操作

集成了在线数据库客户端 CloudBeaver，无需安装，通过可视化界面进行数据库管理和查询，提升了数据操作的易用性和效率。

(6) 全生命周期的 SQL 管控

提供 SQL 全流程的管控视角，帮助您统一管理 SQL 质量。您可以追踪问题 SQL 的解决进度，并提供快捷的优化功能，以提升 SQL 的效率。

总结：

综合当前趋势和方向，我们可以看到 SQL 审核工具不仅仅在传统的安全和性能审核领域发挥着作用，它们正在向集成化、智能化转型。云服务集成和权限管理等高级功能正在成为企业的新需求。尽管面临技术挑战和人才短缺的问题，但是这一领域充满了机遇。企业需要选择能够适应快速变化市场的 SQL 审核工具，同时，工具供应商需要不断创新，以满足不断演进的市场需求。

8、数据复制与集成

一、数据复制与集成的背景

同数据库系统一样，数据复制与集成是一个古老传统的技术领域。数据集成的目的在于将不同来源的数据集成合并，为用户提供统一的数据视图。第一个数据集成系统最早于 1991 年被设计出来，主要用于集成公共微数据。此后，随着 IT 信息化、数字化进程的推进，利用数据复制与集成工具进行跨同异构系统之间的数据复制集成成为各企业应对商业挑战的行业共识。

借助数据复制与集成工具，企业可以快速应对如下业务场景：

1. **数据库迁移**，因数据库大版本升级、上云迁移、下云迁移、跨环境迁移等业务诉求，存在数据库迁移的业务需求。此时，需要借助数据复制与集成工具进行结构与数据的搬迁。
2. **数据容灾**，考虑到业务高可用，很多行业客户会构建业务与数据容灾中心，此时需要借助数据复制与集成工具进行数据的容灾复制。
3. **数据仓库集成**，数据仓库作为企业统一数据存储与分析引擎，需要借助数据复制与集成工具将企业内外部各个数据孤岛的数据统一集成至数据仓库。
4. **数据共享与交换**，在企业之间或企业内部，经常会出现共享或交换数据的需求，此时也需要借助数据复制与集成工具进行数据的离线或实时集成。

二、数据复制与集成的核心技术

伴随着数据复制与集成技术的发展，衍生出了多种不同的集成技术。企业可根据业务的不同特点需求，选择使用不同的集成方案。

1. **批量数据集成**，能够通过 ETL(Extract、Transformation、Load)/ELT(Extract、Load、Transformation)等技术，将数据源中的数据批量或批处理模式进行访问、采集、转换与集成。批量数据集成通常用于大数据、数据仓库的周期性数据入仓集成。对于这种数据集成技术，最核心的能力要求是数据集成的速度与质量。传统的数据集成工具一般都属于这种技术方案，例如：海外老牌数据集成厂商 informatica、kettle、Oracle data integrator 等。

2. **数据复制及同步**，当今商业社会的竞争愈发激烈，企业为了能够在商业竞争中占得先机，重度依赖于实时数据的挖掘与应用。为此，数据集成领域衍生出实时数据复制与同步技术，以提供更实时的数据集成能力。所谓的数据复制与同步是指，能够从数据库、文件、消息队列等数据存储系统中，以近乎实时的速度进行数据的访问、采集和集成。数据复制及同步技术中，最典型的技术为：Log-based change data capture(CDC)。对于数据复制及同步技术，最核心能力在于数据实时采集技术，以及数据复制的延迟。基于数据复制及同步技术，一般可用于满足：数据库迁移、数据容灾、数据多活、数据仓库实时集成等业务场景。当前业界新一代数据集成厂商都在这个技术上提供较强的竞争力，例如：海外数据集成厂商 Fivetran、airbyte，以及国内玖章算术的 NineData。

3. **数据虚拟化**，数据虚拟化是另外一种数据集成技术，其能够基于分析引擎，通过 SQL、JDBC 或 API 等接口方式进行跨同异构数据源的统一数据查询分析。基于数据虚拟化技术，企业不需要进行数据的搬迁，即可基于分布式分析引擎快速进行多源数据库的统一查询分析。这种技术的优势在于，不需要进行前期的数据 ETL/ELT 过程，可以直接进行数据分析使用。但由于，数据存储计算强依赖下游数据源，为此分析性能有一定瓶颈，比较适合应用于低频的数据查询分析。数据虚拟化技术比较有代表性的产品有：Apache trino，以及国内玖章算术 NineData 的 DSQL。

在数据复制与集成过程中，为了更好地满足业务场景、保障集成成功率，还衍生出如下的相关技术：

1. **数据处理**。在数据集成过程中，业务很多时候有进行轻量数据清洗处理的诉求。为此，很多数据集成工具一般都会提供数据处理能力。例如：半结构化数据解析、数据自填充、数据过滤、数据计算、数据脱敏等能力。

2. **质量保障。**在数据集成过程中，数据被进行搬迁移动，很容易出现因软件、硬件问题导致的数据质量问题。伴随着数据集成技术，数据质量保障相关技术也得到相应的发展。质量保障方面主要有如下几种相关技术：

(1) **数据质量检测。**质量特征检测需要用户对数据有一定程度的理解，其能够根据数据情况，定义数据质量特征规则，然后基于这些规则进行数据质量检测。例如，用户可以定义数据中某字段的最大值、最小值、平均值、空值数量、重复值个数等特征规则，基于这些特征验证数据是否存在质量问题。

(2) **全量数据对比。**全量数据对比会进行数据复制的源与目标两个数据源之间对应数据的精准数据对比。其会根据源与目标数据源的不同字段内容进行精准对比，精准找出不一致数据的记录及字段。这种对比方案比较适用于数据复制的两个数据源的数据质量检测。目前业界数据集成工具一般都会配套对应的数据对比能力，例如玖章算术的云原生数据管理 NineData，阿里云数据传输服务 DTS。

(3) **快速数据对比。**快速数据对比，只进行数据复制的源与目标数据源的部分数据的抽样对比，并进行部分特征对比，例如表记录数、某字段的最大值、最小值、平均值等。相较于全量数据对比，快速数据对比的速度更快，对比精准度相对较差。

三、数据复制与集成的发展趋势

数据复制与集成技术虽然已经发展了 30 年之久，但伴随着底层基础 IT 设施的发展，商业社会的进步，其也呈现出新的发展趋势，具体如下：

1. **云原生化。**传统数据复制与集成软件一般以线下软件形式实施交付，这种交付方式效率低，使用门槛高，后期运维挑战大。随着云计算的发展深入，IT 上云已经成为企业的普遍共识，据 Gartner/IDC 等咨询机构的统计分析，企业在 Cloud 的投入已经超过私有数据中心的投入，且 2023 年云数据库的占比也已经超过传统数据库。当然，数据复制与集成领域的云原生化也就成为必然的趋势。基于云计算 IAAS 能力高效构建并提供数据复制与集成云服务，可有效降低数据复制使用门槛、实现更高效的运维保障，提升交付效率。

2. **一体化。**随着企业对数据的使用规模、使用深度的深入，必然存在批量数据集成、实时数据集成等不同集成方式的要求；同时，也会存在丰富多样的孤岛数据源集成诉求。从技术栈统一、降低运维维护成本、减少技术人员技能培养投入等角度考虑，通过统一的数据集成平台实现所有数据源的离线、实时数据集成将是企业的最佳技术选择。为此，

数据集成厂商未来将会逐步扩张自身产品边界，提供更一体化、更统一的数据复制与集成服务。

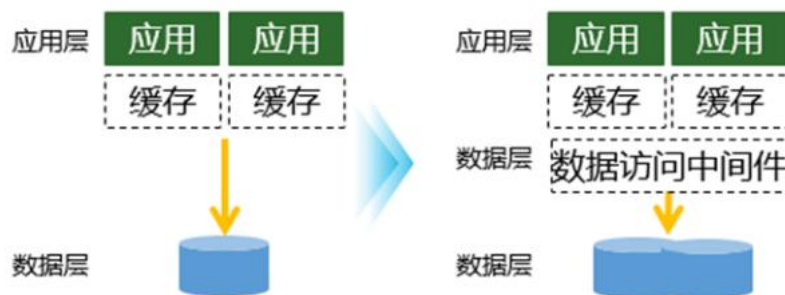
3. **Serverless**。企业业务通常都存在周期性，并非每个时刻的业务量都是相同的。例如互联网应用白天一般是流量高峰期，此时，产生的数据量一般也是高峰期，而晚上很多应用都没有流量，此时需要复制集成的数据量也会相应下降。如果数据复制与集成工具不能根据业务数量进行自适应的资源弹性伸缩，那么业务就需要按照最大业务峰值进行资源的预估与投入，存在极大的资源浪费，且业务一旦有突发流量时，也会出现因预估资源不足导致的数据交付延迟等问题。为此，数据复制与集成工具，需要具备 Serverless 的自动弹性伸缩能力，能够根据业务数据量的变化趋势，进行集成复制资源的自动升缩。

4. **智能化**。除了结构化数据外，据 IDC 预测半结构化/非结构化数据在数据中占比高达 90%。随着 AI 技术的发展，这些数据的价值能够被逐渐挖掘并应用。在数据集成领域，企业想要进行半结构化/非结构化数据的集成，依然存在较大的技术挑战。为此，未来的数据集成工具要能够深度集成 AI 能力，进行元数据的智能分析，并可以自动执行数据的采集、转换和集成任务。

9、数据库中间件

一、需求状态

随着数据规模的不断扩大、数据访问的频率、并发访问量的不断增大，数据库为了更好支撑数据量和访问的需求，开始广泛采用集群部署和分布式架构，比如，对于集中式数据库经常采用主从、读写分离的技术，原生分布式数据库采用分片和复制的技术。相应地对于数据的使用方为了适应数据层的变化，也需要进行架构的调整，从简单的数据库驱动、ODBC/JDBC、连接池管理逐渐发展出更复杂的、适应数据层部署，并能进行相应分布式优化的数据库中间件层。



一般来说，通过数据库中间件层实现：

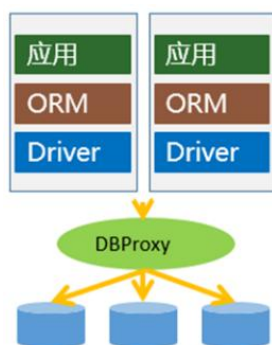
- 为应用屏蔽分布式数据库访问架构上的复杂性，提供更统一的数据库服务编程访问接口
- 通过统一的负载管理、可靠性算法提供更可靠的数据访问能力
- 为分布式查询、特别是 join 操作提供更一致的优化手段
- 提供性能更高的分布式事务管理支持
- 方便提供更全局一致的热点数据缓存
- 提供统一的安全管理机制

二、典型实现方式

基于数据访问层部署方式和形态上的差异，数据库中间件从原理上通常有 2 种实现方式，不同的方式在实现和运维上各有优缺点。

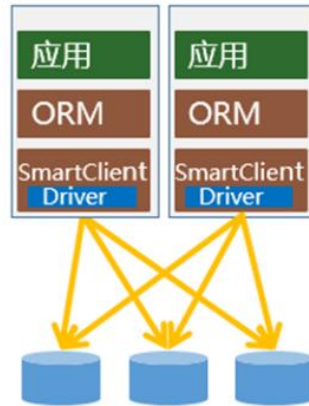
（一）服务端负载均衡方式

从应用中抽象出单独部署的数据库访问服务层以 DBProxy 的方式代理应用对数据的访问。代理服务北向提供应用无感的数据访问接口，南向负责数据库访问的细节，一般包括分布式数据库连接管理，分库分表/分片配置和管理，SQL 解析和改写等。基于 DBProxy 也可以方便地实现统一的分布式数据库管理。



（二）客户负载均衡方式

数据库访问层嵌入到应用中，以分布式数据细节感知的方式(SmartClient)访问数据。数据库访问组件通常以库方式与应用打包部署，通过直接访问数据库访问以获得某些数据库直接访问的性能优势。当然，分布式部署的组件会在维护和配置更新，以及缓存一致性管理方面会产生一些开销。



三、常见产品方案

1. Apache ShardingSphere



ShardingSphere(<https://shardingsphere.apache.org/>)是开源的分布式数据库中间件产品和解决方案组成的生态，它主要由 JDBC、Proxy 和 Sidecar（规划中）这 3 款相互独立，却又能够混合部署配合使用的产品组成。这些产品均提供标准化的数据分片、分布式事务和数据库治理功能，可适用于如 Java 同构、异构语言、云原生等各种多样化的应用场景。其主要特性和优势如下：

- (1) 数据库分片：ShardingSphere 支持垂直分片和水平分片，可以根据业务需求和数据特征将数据分布在多个数据库实例中，从而提高系统的扩展性和吞吐量。
- (2) 分布式事务：ShardingSphere 提供了分布式事务的支持，包括 SAGA 模式和全局事务，可以保证在分布式环境下的数据一致性和 ACID 特性。
- (3) 数据加密：ShardingSphere 支持数据加密功能，可以对敏感数据进行加密存储，以保护数据安全和隐私。
- (4) 弹性扩展：ShardingSphere 允许在线动态添加或移除数据库节点，支持自动的分片键扩展和缩小，使得系统能够根据实际负载进行弹性扩展。

- (5) 透明化：ShardingSphere 对应用程序是透明的，无需修改应用程序代码即可享受到分片、事务和加密等特性，降低应用开发复杂性。
- (6) 易于集成：ShardingSphere 可以轻松集成到现有的应用架构中，支持多种数据源类型，如 MySQL、PostgreSQL、Oracle、SQL Server 等。
- (7) 社区支持：ShardingSphere 拥有活跃的开源社区，提供了丰富的文档、教程和案例，可以帮助用户快速上手和解决实际问题。

2.ProxySQL



ProxySQL(<https://www.proxysql.com/>)是一个高性能、高可用、协议感知的 MySQL 代理，采用 GPL 许可协议，可以免费使用和访问。提供了 MySQL 服务器之间的数据复制和查询路由功能，专为解决大型 MySQL 部署中的可扩展性、高可用性和故障转移问题而设计的，其主要特点和优势如下：

- (1) 负载均衡：ProxySQL 可以根据不同的负载均衡算法（如轮询、最少连接、IP 哈希等）将客户端请求路由到不同的 MySQL 服务器上，从而提高系统的整体性能和吞吐量。
- (2) 故障转移：ProxySQL 支持 MySQL 的主 - 从复制，可以在主服务器发生故障时自动切换到从服务器，确保系统的持续可用性。它还可以检测到从服务器的故障，并将其从路由中排除，以避免死锁。
- (3) 读写分离：ProxySQL 可以将读请求和写请求分开，将读请求路由到从服务器，将写请求路由到主服务器，从而提高系统的读取性能和写入性能。
- (4) 动态路由：ProxySQL 可以根据后端 MySQL 服务器的实时状态动态调整路由策略，例如根据服务器的状态、延迟、响应时间等参数进行路由决策。
- (5) 性能监控：ProxySQL 提供了丰富的性能监控指标，可以帮助管理员实时了解系统的运行状况，并根据需要进行优化和调整。
- (6) 易于管理：ProxySQL 可以通过简单的命令行工具或 Web 界面进行配置和管理，使得系统部署和维护变得简单易行。
- (7) 兼容性：ProxySQL 支持 MySQL 的常见版本，包括 MySQL 5.7、MySQL 8.0 等，并可以与现有的 MySQL 集群和工具进行无缝集成。

四、发展趋势

数据库访问中间件技术发展趋势主要体现在以下几个方面：

- **云原生和容器化**：随着云计算的普及，云原生和容器化技术逐渐成为数据库访问中间件的重要发展方向。这些技术可以帮助企业更好地实现快速部署、动态扩展和高可用性，以满足不断变化的业务需求。
- **微服务架构支持**：随着微服务架构的流行，数据库访问中间件需要支持更多的微服务场景，提供相应的解决方案和工具，帮助企业实现微服务架构下的数据库管理和访问。
- **数据安全和隐私保护**：随着数据安全和隐私保护问题的日益突出，数据库访问中间件需要提供更加完善的安全功能，包括数据加密、身份验证、访问控制等，以确保数据的安全和隐私。
- **优化和监控增强**：基于运行时数据监控，数据库访问中间件可以进行更好的查询解析、处理和查询缓存机制，针对后端的数据库层提供丰富的优化策略。
- **智能化和自动化**：人工智能和机器学习技术的发展为数据库访问中间件的智能化和自动化提供了新的可能。通过这些技术，数据库访问中间件可以自动优化性能、预防故障、智能推荐等，提高系统的稳定性和效率。
- **支持多模数据访问**：随着大数据和物联网技术的发展，企业需要处理的数据类型越来越复杂。数据库访问中间件需要支持更多的数据模型，如关系型、非关系型、时序数据等，以满足企业多样化的数据处理需求。
- **提供数据管理能力**：作为数据库管理工具，通过数据库中间件实现数据迁移和管理能力，数据加密脱密方案。

总之，数据库访问中间件技术的发展趋势是多元化、复杂化和智能化的，需要不断进行技术升级和创新以满足市场需求和企业需求。

总 结

在 2023 年的征途中，中国数据库行业迎来了新的发展里程碑。技术创新继续作为推动行业前行的核心动力，云原生、分布式数据库技术的广泛应用，以及 AI 与数据库融合的智能化趋势，不仅增强了数据库的性能和适应性，也促进了产业链的深度延伸。国产数据库在国际市场上的活跃表现，逐步提升了其全球影响力，通过国际合作和技术交流，中国数据库品牌正在世界舞台上绽放光芒。

生态建设方面，国产数据库厂商通过与合作伙伴的紧密协作，成功培育了一批活跃的技术社区和开源项目，为行业的可持续发展奠定了坚实的基础。资本市场对国产数据库领域虽有所降温，但市场对创新技术的追求并未减少，依然有优秀的初创企业获得资金支持，显示出投资者对国产数据库未来发展潜力的信心。

开源生态的发展同样值得关注，国产数据库的开源化进程加速，不仅推动了技术共享和创新，也为国产技术的国际化发展打开了新的大门。在金融核心系统的国产化进程中，国产数据库展现了其在关键应用场景中的稳定性和可靠性，一系列成功的案例证明了国产数据库已经具备了承载重要业务系统的能力。

国家政策的支持和市场需求双重驱动下，国产数据库行业发展迎来了新的机遇。未来，我们期待国产数据库厂商能够继续加强技术创新，深化行业合作，提升产品质量和服务水平，共同推动中国数据库行业的繁荣发展。随着数字化转型的不断深入和新兴技术的持续涌现，中国数据库行业有望在全球舞台上发挥更加重要的作用，为全球数字经济的发展贡献中国的力量。