

Real-Time Computing in a Dynamic, Open Environment: The Broad Learning Approach for Offline Intelligent

C. L. Philip Chen

Dean, School of Computer Science and Engineering,

South China University of Technology,

Guangzhou, Guangdong, China

Co-President, Guangdong AI Industry Association

Former Vice-President, Chinese Automation Association

Former Vice-President, Macau S&T Association

Former President, IEEE SMC Society

Director, **Ministry of Education**, Health Intelligence and
Parallel Digital Human Engineering Research Center

Director, **Guangdong** Computational Intelligence Modelling
and Perception **Key Lab**

Vice Director, **Guangdong** AI and Digital Economy Lab

Honor President, **Shenzhen** AI Industry Association

National Distinguished Expert

IEEE Life Fellow, AAAS Fellow, IAPR Fellow, CAA, CAAI,
HKIE Fellow

Member of Academia Europaea (M-AE)

Member of European Academy of Science and Arts (EASA)

Foreign member of Russian Academy of Engineering (RAE)

Philip.Chen@ieee.org



C. L. Philip Chen



Member of Academia Europaea
Member of European Academy of
Science and Arts (EASA)

1989-2001
Wright State
University,
Ohio

2002
University of
Texas, San
Antonio

2007 IEEE
Fellow/AAAS
Fellow/IAPR
Fellow, 2015

2012-2013,
IEEE SMC
Society

2014--2019 -2021
IEEE Trans on SMC:
Systems/Cyberneti
cs Editor-in-Chief
(IF 13.451/11.488)

2019-, Dean,
School of
Computer Science
and Eng

2021, Wu WJ AI
Distinguished
Contribution Award,
2018, 2021年, IEEE
Norbert Wiener,
IEEE Joseph G. Wohl
Award

1985 MS
University of
Michigan

1988 Ph.D.
Purdue
University

Visiting Research
Scientist, 1990-2000,
US WP AFB, Wright
Lab, NASA Glenn
Research Center

2005 UTSA,
Dept head and
Associate Dean

2010, Dean of
Faculty of S&T,
University of
Macau

2018 Elected
member of
EA&EASA

2018-
2026
HCR

- IEEE Fellow, AAAS Fellow, IAPR Fellow, 中国自动化学会Fellow
- 中国自动化学会副理事长;
- 国际总主席, IEEE系统人机及智控学会(SMCS) (2012-2013);
- IEEE SMCS Fellow 评选主席 (2016-2017)、审核委员 (2009-2018);
- 现任CAAI Trans. on AI执行主编, 曾任IEEE Trans. Cybernetics和IEEE Trans. on SMC: Systems期刊主编及中国科学等多个期刊副编, 自动化学报副主编;
- 2018和2021两次获IEEE TNNLS最佳论文奖;
- ESI前1%高被引论文58篇、3篇入选ESI前0.1%热门论文;

- SCI期刊文章1000余篇(其中, IEEE Transactions期刊文章600余篇), 学术专著1部, 授权美国专利4项;
- 2018年-2024年连续6年入选科瑞唯安全全球高被引科学家; 5次IEEE SMC学会杰出贡献奖得主;
- 主持科技部重点研发计划、973课题, 国家自然科学基金重大研究计划培育项目等10余项基金;
- 2018年获 IEEE系统人机控制论的最高学术奖--IEEE 诺伯特·维纳奖 (Norbert Wiener Award)
- 2021年度IEEE Joseph G. Wohl (约瑟夫·沃尔) 终身成就奖
- 荣获2021年度吴文俊人工智能杰出贡献奖



Guangzhou: capital of Guangdong Province



GUANGZHOU INTERNATIONAL CAMPUS
SOUTH CHINA UNIVERSITY OF TECHNOLOGY

Guangzhou China's "Southern Gateway"

A window connecting China and the world

Population: 19.1 million residents

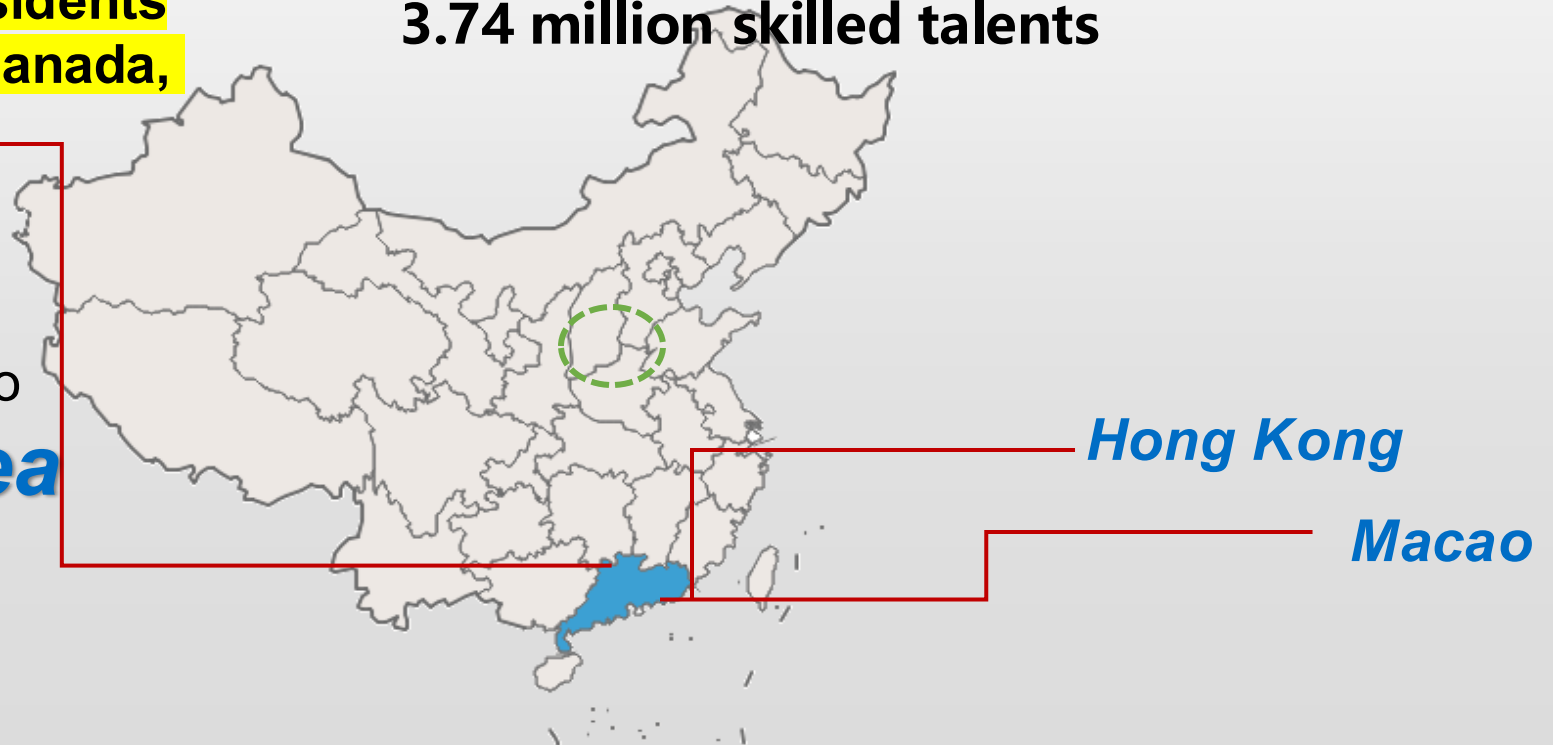
Guangdong Province

- ✓ Population: 128.1 million residents
- ✓ GDP: 2.05 Trillion equiv to Canada, and Russia, and Brazil

- The world's only major trading port that has remained prosperous for over 2,000 years
- One of the fastest-growing megacities in the world
- 13,000 high-tech enterprises, 80 new R&D institutions, 1.65 million university students, and 3.74 million skilled talents

Guangdong, Hong Kong & Macao

Greater Bay Area



South China University of Technology

All three campuses are in Guangzhou



Wushan Campus

- 1.83 millionm²
- 28,000 students
- Engineering



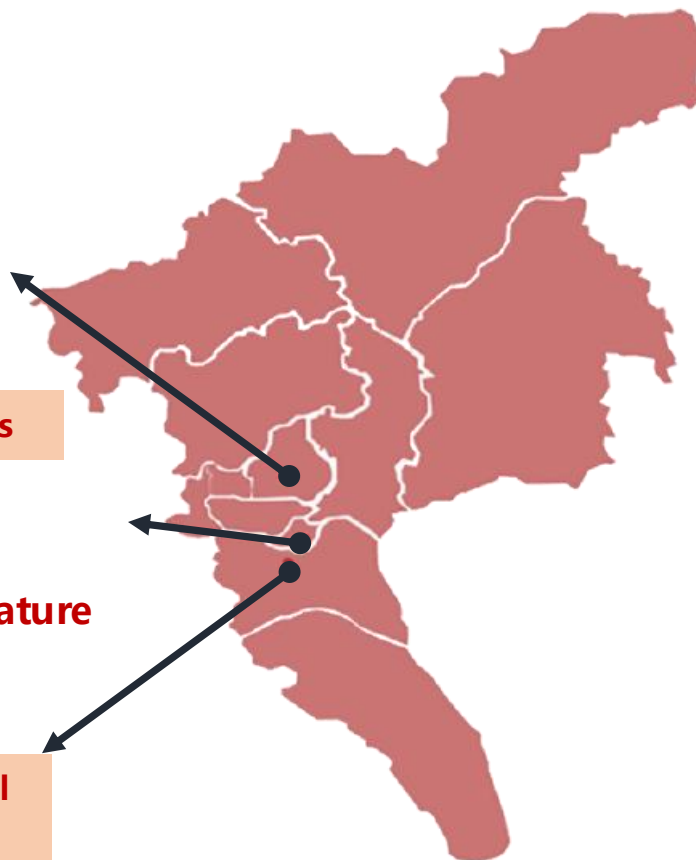
University Town Campus

- 1.11 millionm²
- 17,000 students
- Information & Literature Arts



Guangzhou International Campus

- 1.13 millionm²
- 10,000 students
- Cross-disciplinaries



Campus size **4.47 millionm²**
Building size **2.84million m²**



Fixed assets totaled **10.433 billion yuan**



4,500+ faculty and staff members



55,000+ students enrolled



37 colleges and schools



South China University of Technology (SCUT)

ShanghaiRanks
World University

World
Top 150

2025 Global University
Google Academic
Impact Ranking

Mainland
universities
22nd place

2024 Times Higher
Education Asia
University Rankings

Asian
universities
45th place Mainland
universities
17th place

Nature index ranking

World
Top 50

Ranking agencies		2019年	2020年	2021年	2022年	2023年	2024年
SH Rank		201-300	151-200	151-200	151-200	101-150	101-150
US NEWS		336	290	250	219	219	187
QS		480	462	407	406	392	385
THE		501-600	401-500	401-500	401-500	251-300	251-300

(There are about 45,000 universities worldwide)

Engineering ranks **22nd** globally



20
Disciplines ESI
top 1% globally

6
Disciplines ESI
top 1‰ globally

3
Disciplines ESI
top 1‰‰ globally

2

- Food Science and Technology
- Polymer science

2024 US News World University Subject Rankings
Top three in the world

World-class subject rankings



2 Disciplines Top 10
18 Disciplines Top 50

World University Academic Ranking



3 Disciplines Top 10
18 Disciplines Top 50

Institutions			层次	排名	学校名称	得分	Top Sci Cites	Cites/Paper	Top Papers
1	CHINESE ACADEMY OF SCIENCES		A+	1	清华大学	73.0	326716	16.99	299
2	TSINGHUA UNIVERSITY	30	A+	2	浙江大学	72.8	GEORGIA		USA
3	UNIVERSITY OF ELECTRONIC SCIENCE & TECHNOLOGY	31	A+	3	哈尔滨工业大学	71.6			CHINA MAINLAND
4	SOUTHEAST UNIVERSITY - CHINA	32	A+	4	北京航空航天大学	68.9	TH WALES SYDNEY		AUSTRALIA
5	XIDIAN UNIVERSITY	33	A+	5	上海交通大学	68.1	OF FLORIDA		USA
6	UNIVERSITY OF CALIFORNIA SYSTEM	34	A+	6	北京大学	67.9	CHNICAL UNIVERSITY		CHINA MAINLAND
7	NANYANG TECHNOLOGICAL UNIVERSITY	35	A+	7	华中科技大学	67.8	C UNIVERSITY		HONG KONG
8	UNIVERSITY OF LONDON	36	A+	8	南京大学	67.7			CHINA MAINLAND
9	CENTRE NATIONAL DE LA RECHERCHE SCIENTIFIQUE	37	A+	8	中国科学技术大学	67.7			CHINA MAINLAND
10	ZHEJIANG UNIVERSITY	38	A+	10	北京理工大学	67.6			IRAN
11	BEIJING UNIVERSITY OF POSTS & TELECOMMUNICATIONS	39	A+	10	西安电子科技大学	67.6			USA
12	SHANGHAI JIAO TONG UNIVERSITY	40	A+	12	西北工业大学	66.4	TY - CHINA		CHINA MAINLAND
13	INDIAN INSTITUTE OF TECHNOLOGY SYSTEM (IIT SYSTEM)	41	A+	13	东南大学	63.7	TECHNOLOGY OF CHINA, CAS		CHINA MAINLAND
14	HUAZHONG UNIVERSITY OF SCIENCE & TECHNOLOGY	42	A+	14	同济大学	63.5	ITY		SAUDI ARABIA
15	NATIONAL INSTITUTE OF TECHNOLOGY (NIT SYSTEM)	43	A+	15	电子科技大学	62.6			CHINA MAINLAND
16	NATIONAL UNIVERSITY OF SINGAPORE	44	A+	16	西安交通大学	60.9	ITY		CHINA MAINLAND
17	HARBIN INSTITUTE OF TECHNOLOGY	45	A+	17	华南理工大学	60.7			CHINA MAINLAND
18	UNIVERSITY OF TEXAS SYSTEM	46	A+	18	湖南大学	60.1	ON		ENGLAND
19	WUHAN UNIVERSITY	47	A+	19	中国人民大学	60.0	WEALTH SYSTEM OF HIGHER EDUCATION		USA
20	BEIHANG UNIVERSITY	48	A+						N/A
21	UNIVERSITY OF TECHNOLOGY SYDNEY	49	A+						CHINA MAINLAND
22	EGYPTIAN KNOWLEDGE BANK (EKB)	50	A+						CHINA MAINLAND
23	UNIVERSITY OF CHINESE ACADEMY OF SCIENCES, CHINA	51	A+						CHINA MAINLAND
24	ALPHABET INC.	52	A+				DEFENSE TECHNOLOGY - CHINA		CHINA MAINLAND
25	DALIAN UNIVERSITY OF TECHNOLOGY	53	A+				OSTS & TELECOMMUNICATIONS		CHINA MAINLAND
26	BEIJING INSTITUTE OF TECHNOLOGY	54	A+				OF SCIENCE & TECHNOLOGY		HONG KONG
27	CITY UNIVERSITY OF HONG KONG	55	A+				FORMATION SCIENCE & TECHNOLOGY		CHINA MAINLAND
28	CENTRAL SOUTH UNIVERSITY	56	A+						AUSTRALIA
29	SOUTH CHINA UNIVERSITY OF TECHNOLOGY	57	A+				COLUMBIA		CANADA

LLMs vs Small Models

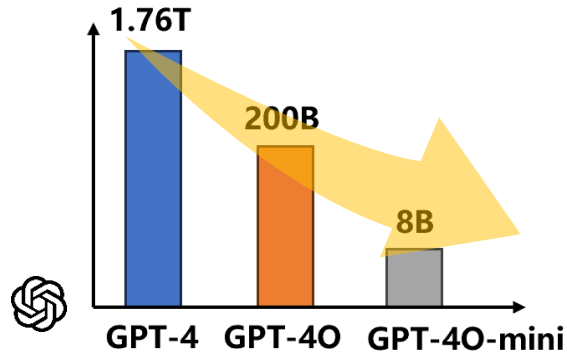
[1] Abacha, Asma Ben, Wen-wai Yim, Yujuan Fu, Zhaoyi Sun, Meliha Yetisgen, Fei Xia and Thomas Lin. "MEDEC: A Benchmark for Medical Error Detection and Correction in Clinical Notes." ArXiv abs/2412.19260 (2024).

LLM has rapidly become a research hot spots, its application potential and industrial driving capacity have strong advantages. However, **it faces challenges** such as high computing power consumption and high deployment costs. **Lightweight large and small models are crucial for reducing computational costs and increasing computing speed.**

Lightweight large models

Compressed by large models through distillation, pruning, and other methods

- Parameter count **Less than one billion**



- Reduces computational resource consumption and increases reasoning speed
- It possesses strong reasoning abilities but relies on primitive LLM making the technology incapable of autonomous control

Pre-trained large language models

Parameter volumes ranging from tens of billions to over 100 billion



GPT-4: 1.76T



Claude 3.5: 175B



DeepSeek V3: 670B



Qwen 1.5: 110B

- Possesses strong reasoning and generation abilities, and has the ability to "emerge."

Autonomous training of small models

The number of parameters can be less than **ten million**, **Lowest computational costs** and **fastest inference speed**

- Task-specific training, Lack of versatility
- YOLO、MobileNet、EfficientNet、Broad learning system
- In the era of LLM, research related to small models is relatively limited



The concept of edge AI agents

Edge AI agents It is an intelligent entity deployed in edge computing environments (usually a combination of software and hardware)., posses **Perception**、**Decision-making**、**Execution** ability。It is not only the "functional unit" of edge computing but also the core of intelligent applications deployed at the edge.

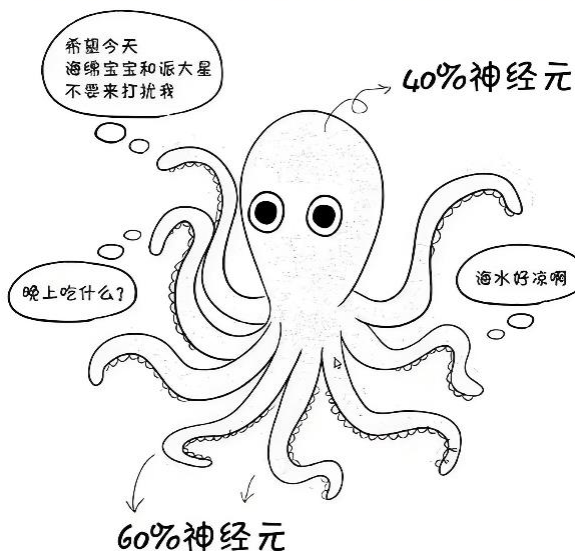
- **Autonomy**: 在本地完成数据处理与决策, 不依赖中心云的实时反馈。
- **Intelligence**: 集成机器学习模型, 支持模式识别、预测、优化等任务。

- **Real-time capability**: 在边缘端快速响应, 满足低延迟需求。
- **Synergy**: 多个边缘智能体之间可形成分布式协作网络, 与云端形成“云-边-端”联动。



The structure of edge agents

章鱼是用“腿”来思考并就地解决问题的



章鱼就是用“边缘计算”来解决实际问题的。作为无脊椎动物中智商最高的一种动物, 章鱼拥有巨量的神经元, 但有高达60%分布在其八条灵活自如的触手(腕足)上, 脑部却仅有40%。

边缘智能体架构

Perception layer:

传感器、摄像头、IoT设备等, 负责采集环境数据。

Edge Computing Layer (Agent Core):

轻量化模型、推理引擎、缓存与本地存储。新型边缘计算具备实时决策、自主学习、隐私保护等能力。

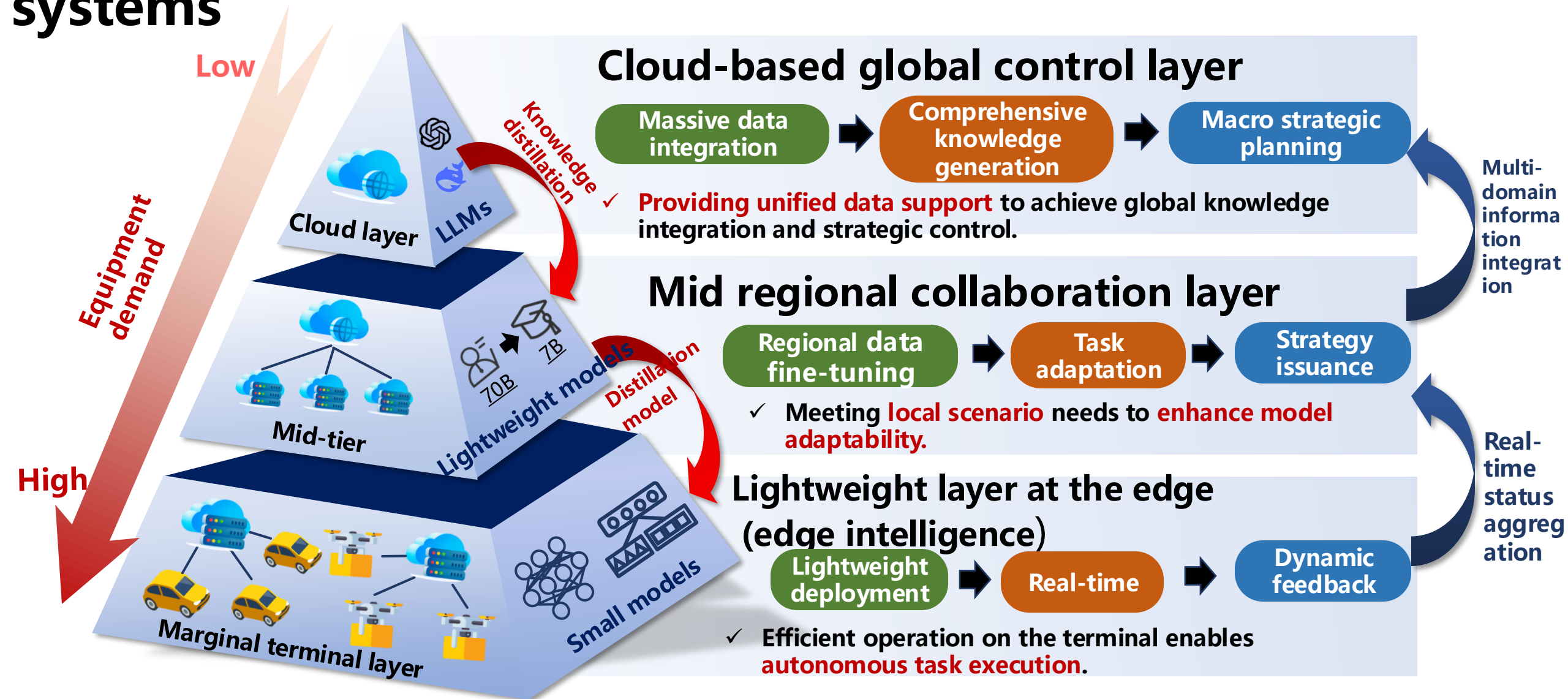
Synergy Layer:

智能体之间的横向协作(边-边), 云端的纵向协作(边-云)

Execution Layer:

将智能体的决策转化为具体动作, 如设备控制、告警、优化指令。

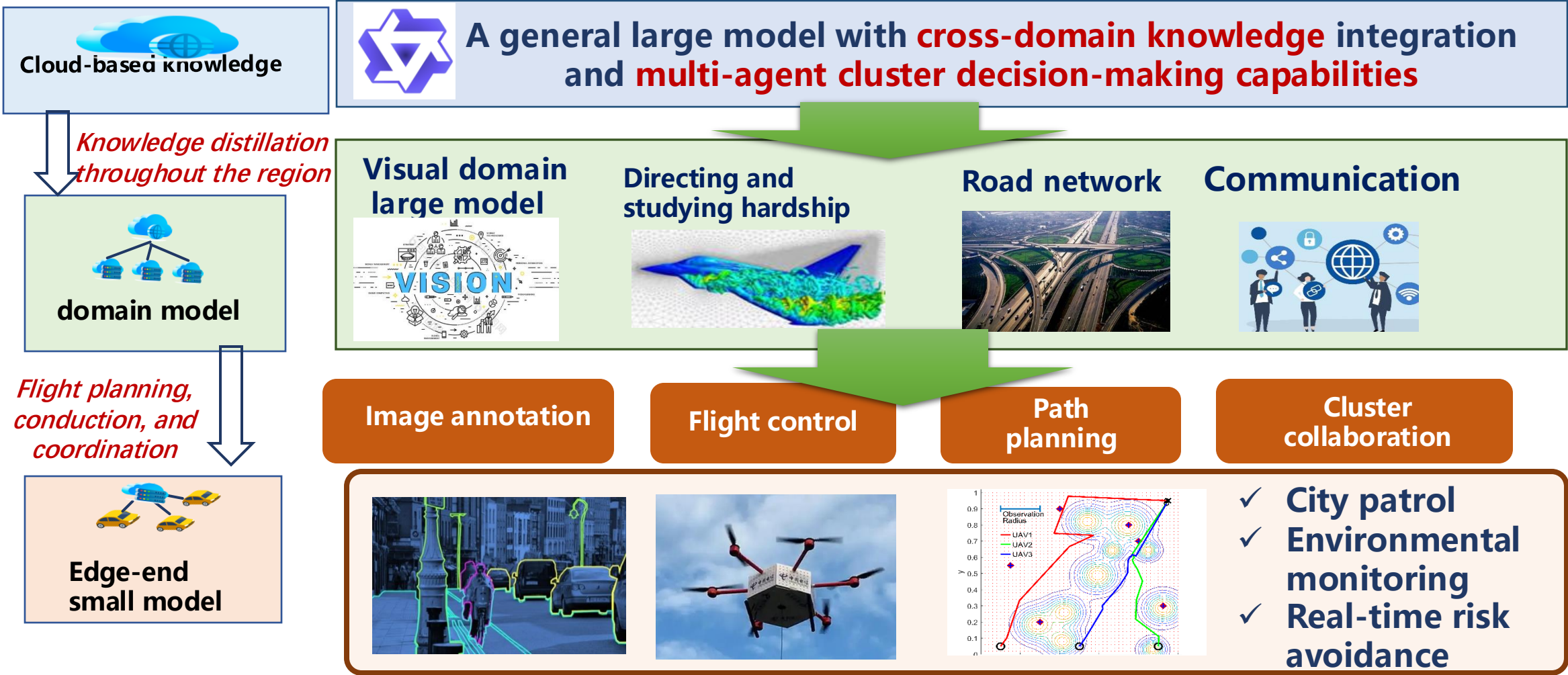
Collaborative large and small models: Cloud-edge AI systems



Build **hierarchical, decoupled large-size** models collaborating with AI systems

Industrialization application scenarios: Drone swarms

Integrating adaptive flight control and dynamic strategic planning,
Enables **efficient and stable cluster-based task execution**.



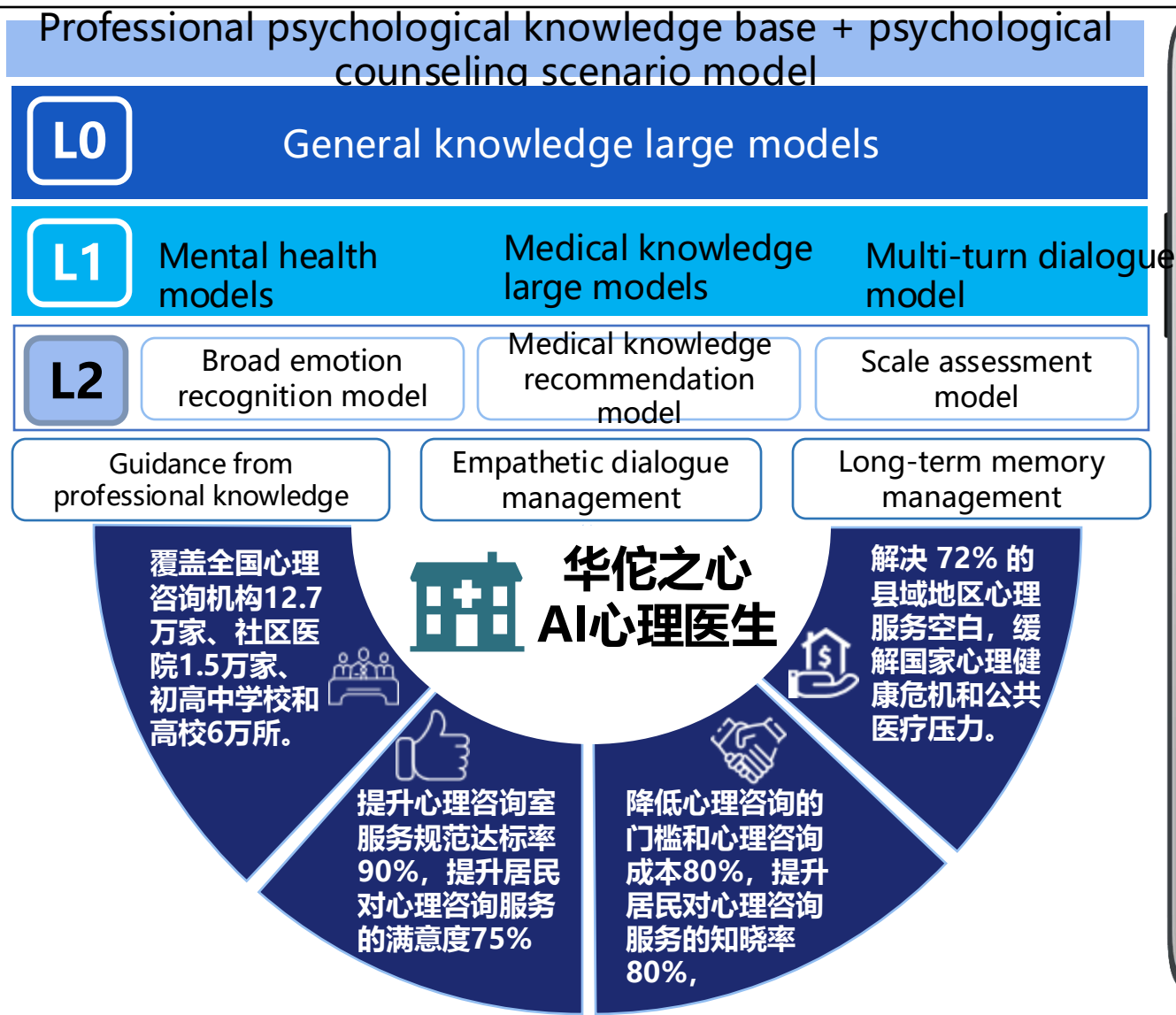
AI large models assist community school psychologist:, Digital intelligence empowers grassroots mental health (网信办立案: 440105230362901260017)

- **华佗之心**项目以心理大模型为技术核心, 以多轮对话的形式精准洞察并响应用户的心理健康需求, 衍生一系列心理健康测评系统、疗愈倾听平台、心理知识科普课堂等心理健康服务产品, 全面实现心理咨询服务的数字化与多元化。

- **华银康-华佗之心**
华银康海珠社区体检中心



- 25年实现广州多个社区试点。26年计划多省推广。
- 按照**专业评测项目付费、疗愈超时长付费、个性化服务等收费方式**向社区、学校、心理咨询机构收取AI心理咨询服务费。



模型、数据、硬件入口

三层技术闭环，支撑 YunovaOS 持续进化。

智能感知平台

脑电、肌电、眼电等感知智能穿戴。

YunovaOS系统

模型推理、云端分析。

类脑智能模型

情绪模型、精神模型、认知模型、意图模

YunovaOS系统，
连接多源信号处理

建模、推理、反馈，形成产品级数据飞轮。

情绪模型

精神模型

认知模型

意图模型

趣乐AI大模型
成功通过国家大模型

国家大模型备案成

云脑智能「趣乐AI」是
全国唯一同时具备脑机接口
神经调控、AI音乐疗愈三赛
合规备案大模型。

备案实力

全国累计868款备案生成式大
专项AI音乐垂类自研模型仅5
趣乐AI为国内唯一的神经调控

技术闭环

依托双模态脑电解码技术，实现
神经反馈干预、个性化疗愈音乐
全链路闭环。

应用价值



脑健康



情绪调节

趣乐AI

国内唯一神经调控音乐大模型

成功通过国家大模型备案

打通 脑机接口 + 神经调控 + AI音乐疗愈

01

五音音乐生成

面向身心平衡与
传统五音疗愈场景，
生成个性化音乐内容

宫 商 角 徵 羽



02

冥想音乐生成

用于放松减压、
专注训练与情绪舒缓

03

睡眠音乐生成

用于入睡辅助、
睡眠陪伴与夜间放松

技术优势

双模态脑电解码
多维解析脑电信号，
精准识别身心状态神经反馈干预
实时神经反馈，
调节脑状态个性化音乐生成
AI生成专属疗愈音乐，
匹配个体需求软硬一体合规布局
软硬一体系统方案，
实现合规落地

未来将持续拓展更多音乐疗愈场景



图1 平陆运河途径地示意图



图2 平陆运河抬填造地与耕地复垦

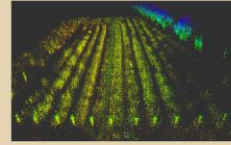
项目名称

主要内容

面向丘陵山区的四足农业智能巡检机器人

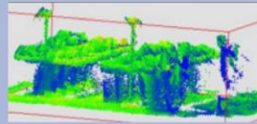
一、多模态融合的丘陵山区混合地图构建与自主导航方法

- ①基于Cartographer框架的三维点云地图构建
- ②基于点云地图的路径规划与精准导航技术



二、基于三维点云的丘陵山区作物表型无损提取模型

- ①基于激光雷达三维点云的作物株高测量方法
- ②基于视觉AI的作物诊断方法



三、动态扰动下基于手势的遥操作鲁棒跟踪控制

- ①基于Leap Motion的手势遥操作控制
- ②动态扰动下四足机器人的鲁棒跟踪控制



现阶段成效

系统集成

作业机器人



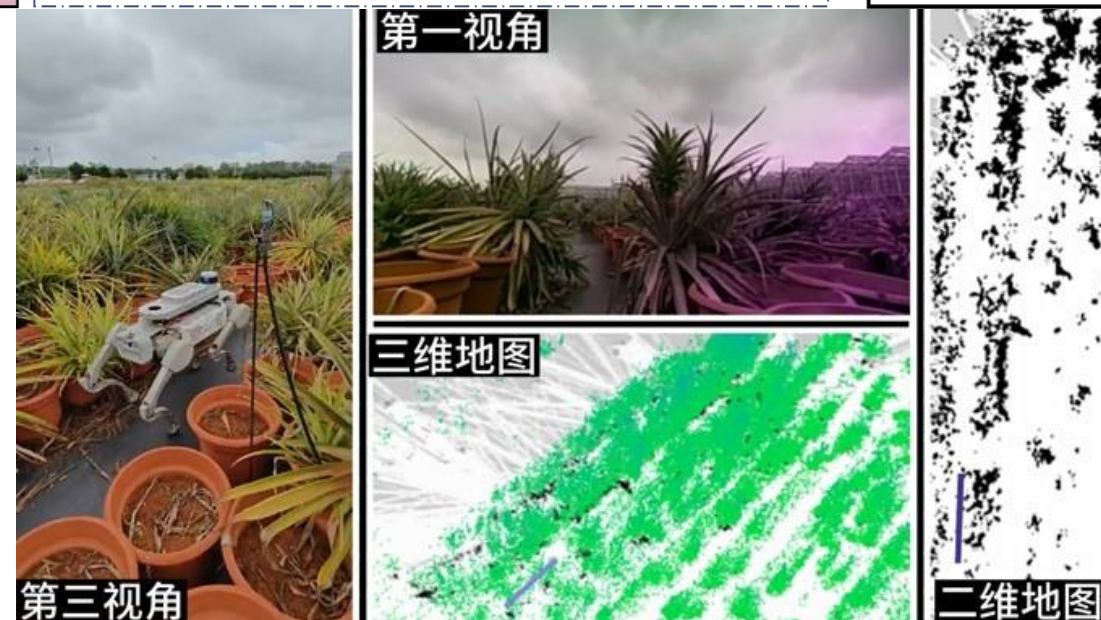
巡检软件平台



爬坡角度 $\geq 30^\circ$
爬梯高度 $\geq 20\text{cm}$

不间断运行4h以上
运动负载 $\geq 20\text{kg}$

诊断准确率 $\geq 90\%$
建图定位误差 $\leq 6\text{cm}$



核心产品



AI servers

AI中心侧推理、
模型训练



AI workstation

边缘侧AI场景
高性能设备



AI accelerator card

搭载华为昇腾
AI芯片平台

云-边-端全场景概念架构



Real-time computation in Edge Embodied Intelligence via: Broad Learning Systems

**Wu WenJun AI Outstanding
Contribution Award, China**

Engineering and Scientific Applications

Optimization, Regression, and Modeling

Least Square Regression

$$\arg \min_{\mathbf{W}}: \|\mathbf{Y} - \mathbf{A}\mathbf{W}\|_2^2 + \lambda \|\mathbf{W}\|_2^2$$
$$\mathbf{W} = (\mathbf{A}^T \mathbf{A} + \lambda \mathbf{I})^{-1} \mathbf{A}^T \mathbf{Y}$$

Least Square Regression

$$\arg \min_W: \|Y - AW\|_2^2 + \lambda \|W\|_2^2$$

$$W = (A^T A + \lambda I)^{-1} A^T Y$$

Almost in real systems, there exist some levels of non-linearity and uncertainty

Deep Neural Networks approach

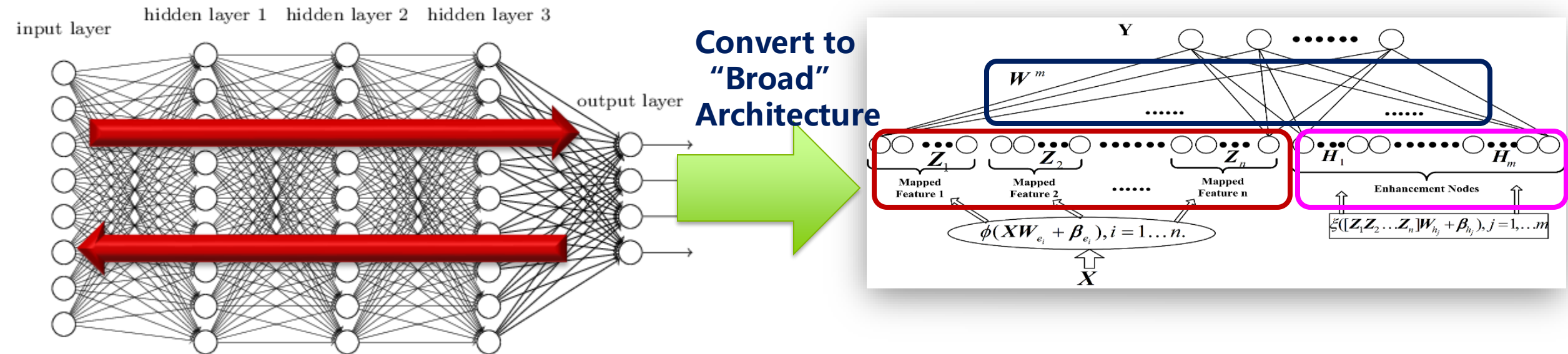
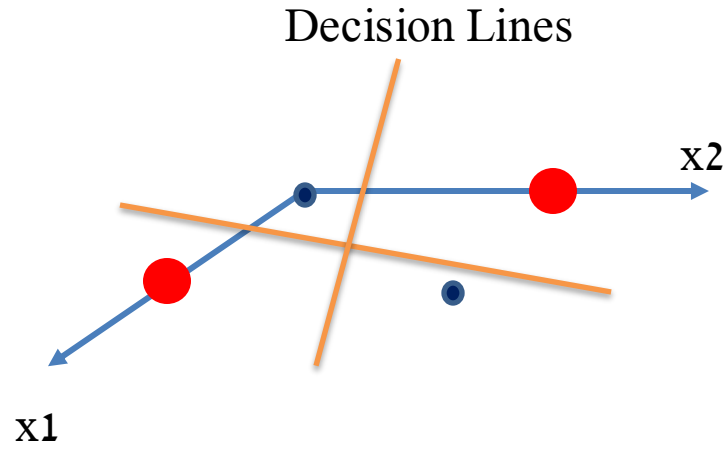


Table 3.4 - Two-Bit Parity in Two and Three Dimensional Feature Space

x_1	x_2	2-bit Parity
0	0	0
0	1	1
1	0	1
1	1	0

Figure 3.15 - Linear Separation of Two Bit Parity Data in Three Dimensions

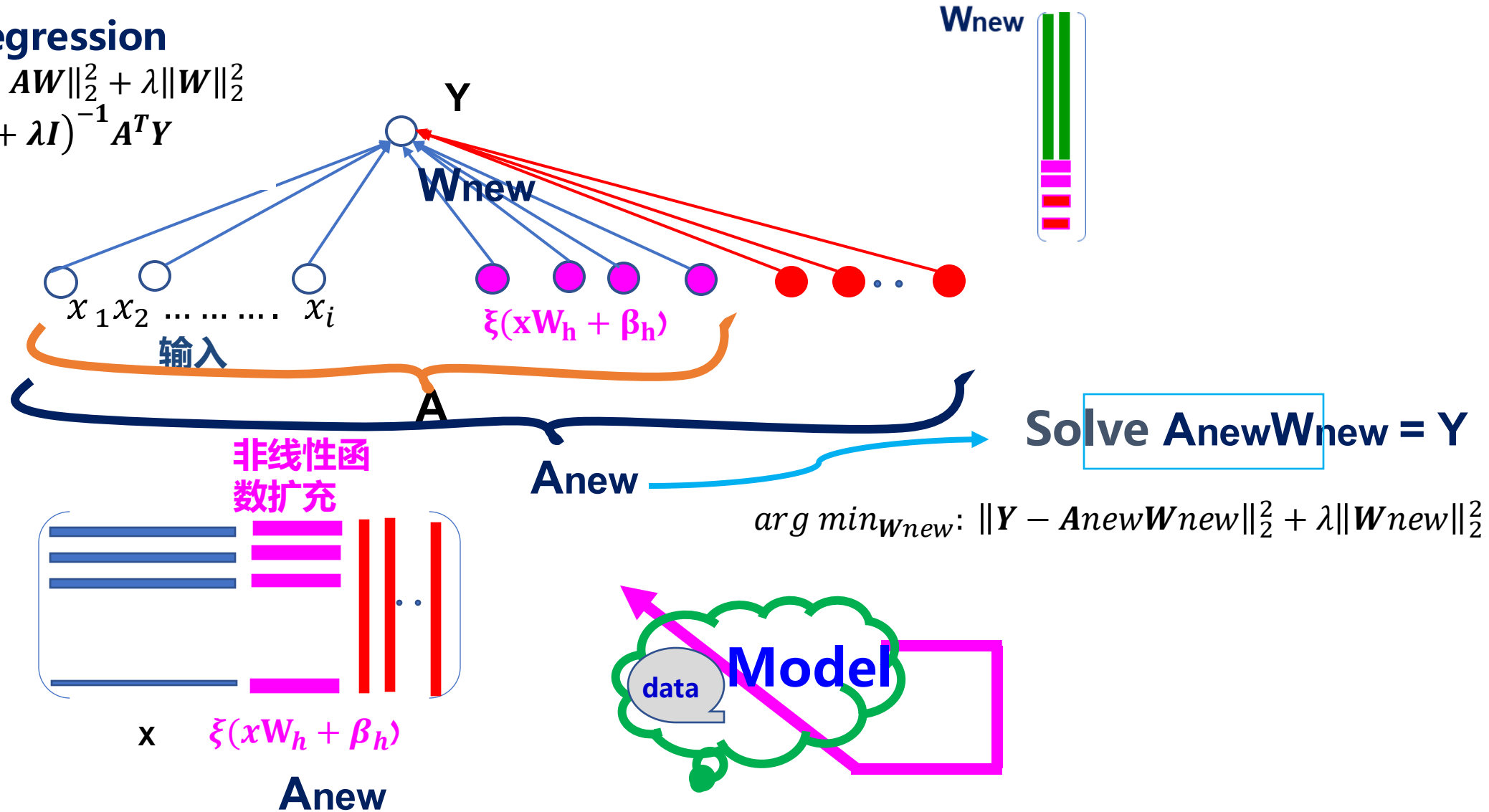


Learning—Reaching the accuracy: Dynamically adding neurons (expanding the structure) --incremental learning

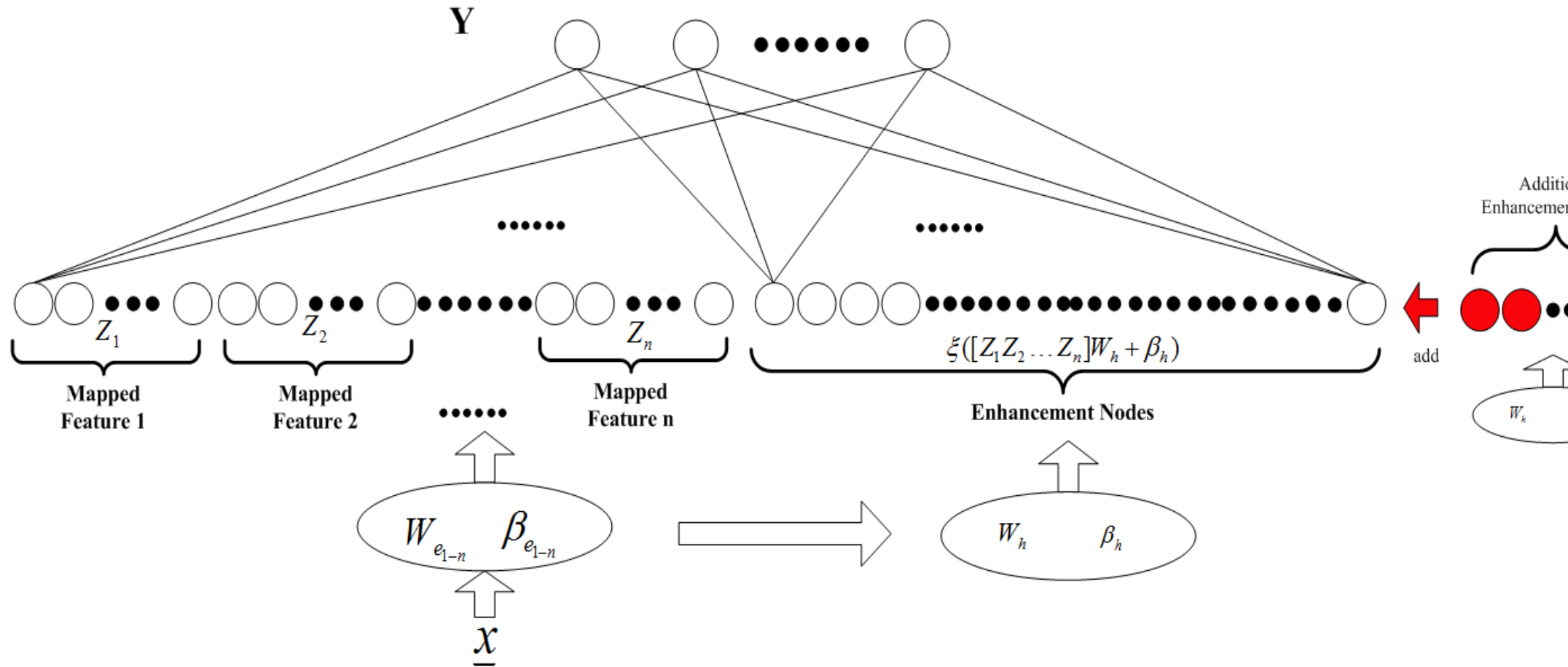
Least Square Regression

$$arg \min_{\mathbf{W}}: \|\mathbf{Y} - \mathbf{A}\mathbf{W}\|_2^2 + \lambda \|\mathbf{W}\|_2^2$$

$$W = (A^T A + \lambda I)^{-1} A^T Y$$



Picturing: Dynamically Adding Neurons: reach learning accuracy

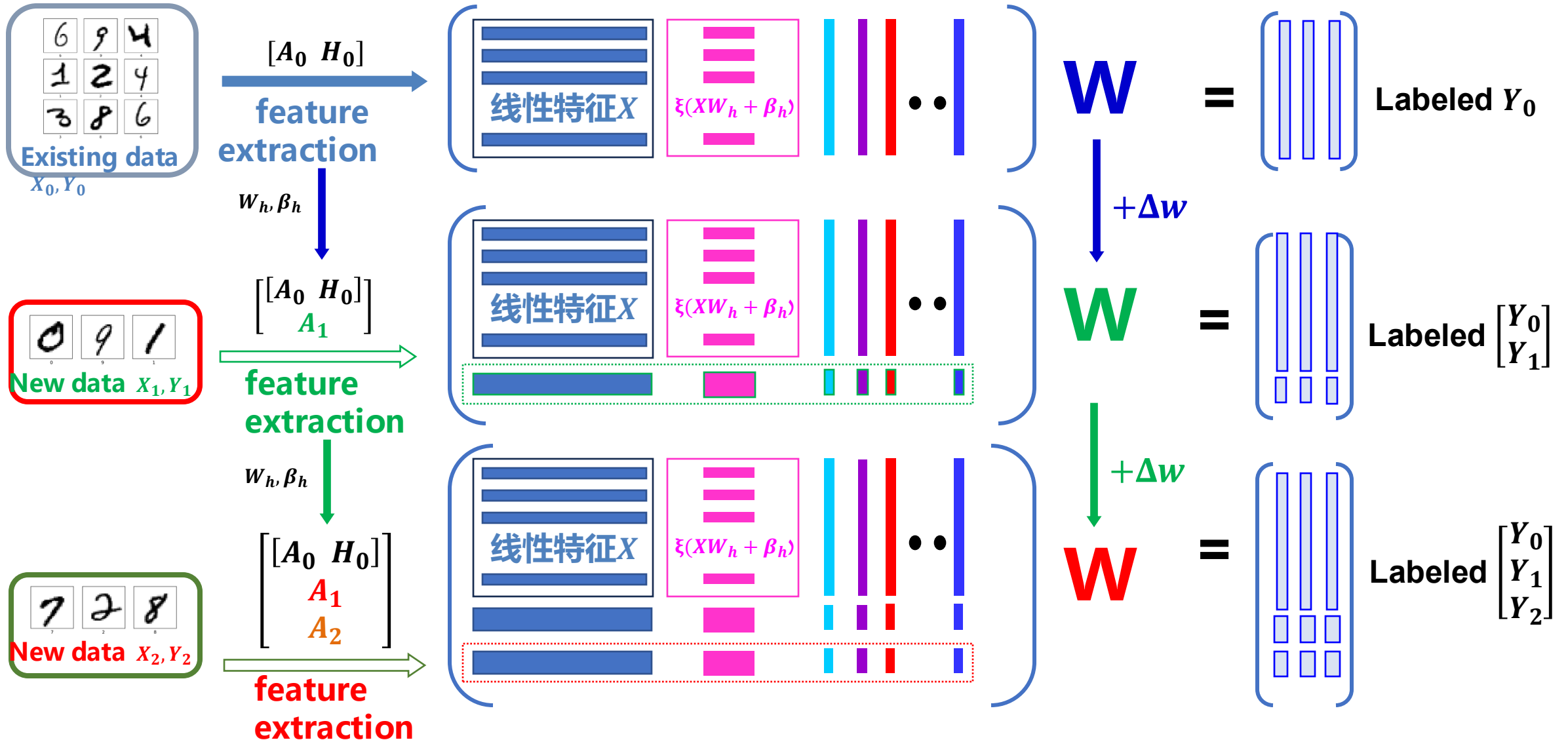


- **CLP Chen** and ZL Liu, "Broad Learning System: An Effective and Efficient Incremental Learning System Without the Need for Deep Architecture," *IEEE Trans on Neural Networks and Learning Systems*, Vol, 29, No.1, pp. 10-24, **2018**.
- **CLP Chen**, ZL Liu, and S. Feng, "Universal approximation capability of broad learning system and its structural variations," 2018 (SMC), , " *IEEE Trans on Neural Networks and Learning Systems*, Vol, 30, No.4, pp. 1191-12044, **2019**.



中国科协年会
Annual Meeting of the China Association for
Science and Technology

Picturing: incrementally learning for new data stream: dynamically restructure the model



Broad Learning System

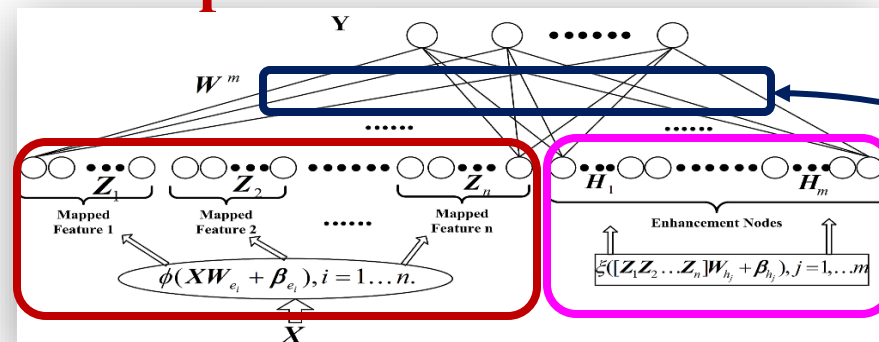


Invented by Philip Chen



Core Theory of BLS and Method: convert Regression to light weighted broad structure and can be computed in real-time

Machine learning is to solve Inhomogeneous non-linear equations $AW=Y$ (i.e., BLS structure), Do not need deep structure learning



Core Modules in BLS

➤ Random Sparse Autoencoder

$$\arg \min_{\hat{W}}: \|Z\hat{W} - X\|_2^2 + \sigma \|\hat{W}\|_1 \quad (1)$$

➤ Random Projection

$$H_j = \xi_j (Z^n W_{h_j} + \beta_{h_j}), j = 1, 2, \dots, m, \quad (2)$$

➤ Feature Extension

$$A = [Z^n, H^m]$$

➤ Least Square Regression

$$\arg \min_W: \|Y - AW\|_2^2 + \lambda \|W\|_2^2 \quad (3)$$

$$W = (A^T A + \lambda I)^{-1} A^T Y \quad (4)$$

Or by SGD Methods

Or by Greville's Method

$$\begin{cases} B = \begin{cases} C^+, & \text{if } C \neq 0, \\ A^+ D (1 + D^T D)^{-1}, & \text{if } C = 0, \end{cases} \\ C = \tilde{A} - D^T A, \quad D^T = \tilde{A} A^+ \end{cases}$$

Finally, the updated weights $\hat{W}$ are analytically computed as:

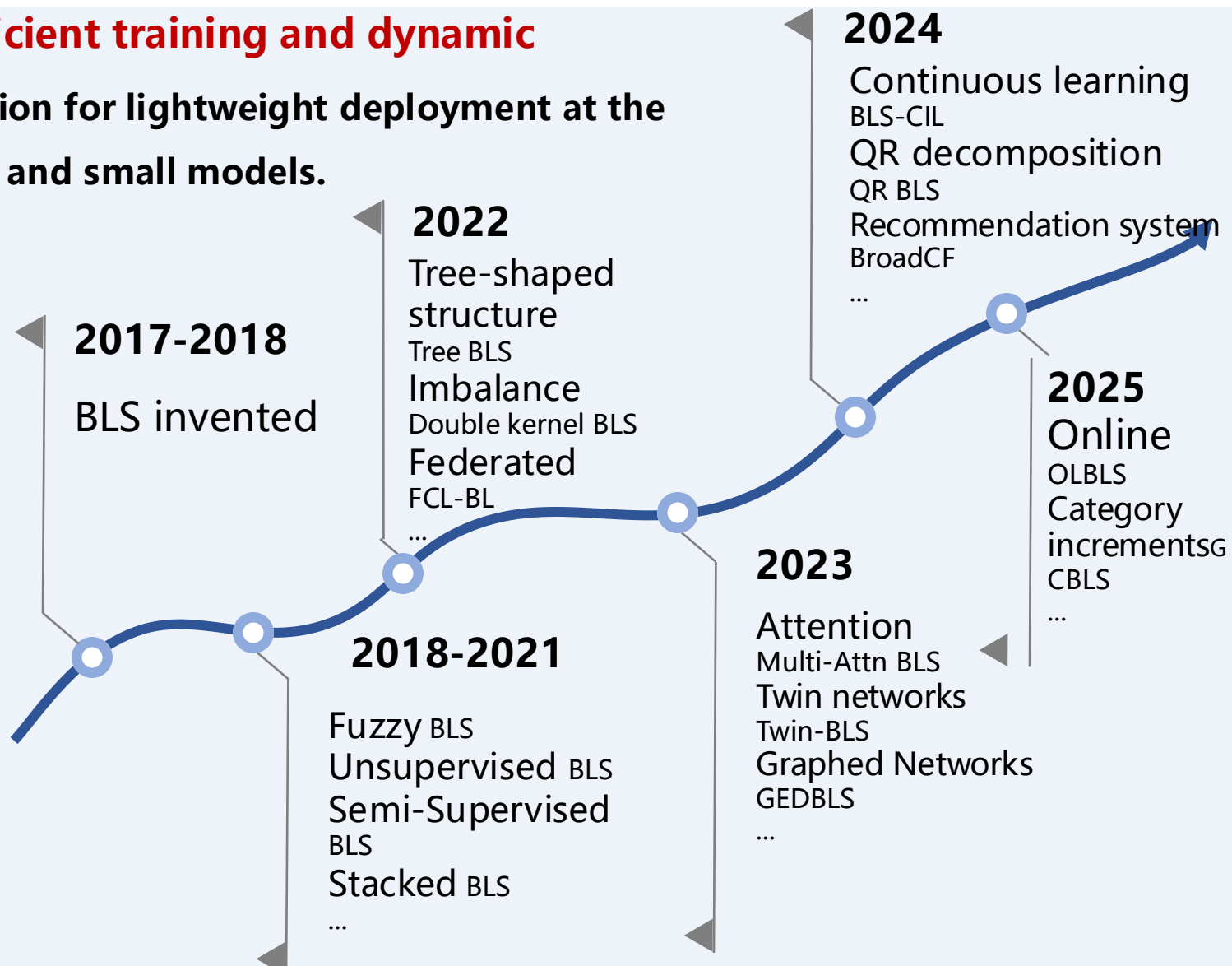
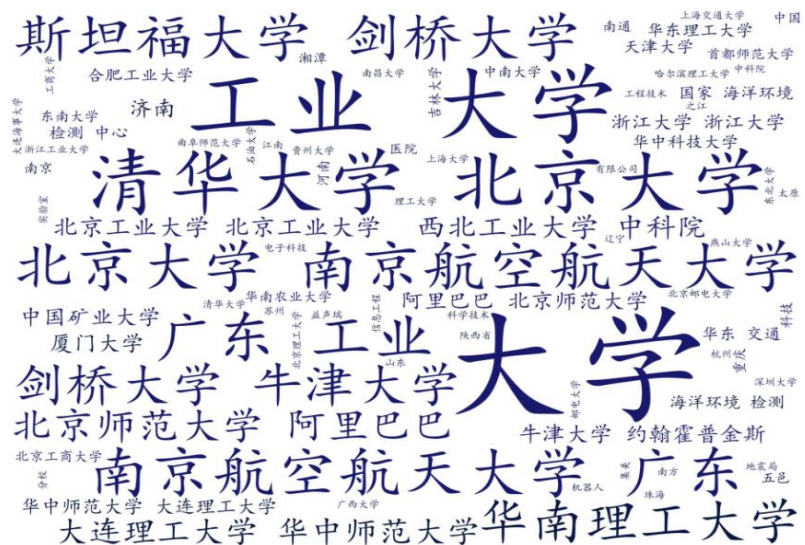
$$\hat{W} = W + B(\tilde{Y} - \tilde{A}W).$$

- CLP Chen and ZL Liu, "Broad Learning System: An Effective and Efficient Incremental Learning System Without the Need for Deep Architecture," **IEEE Trans on Neural Networks and Learning Systems**, Vol, 29, No.1, pp. 10-24, 2018. (影响因子: 11.683, Top期刊, 高被引论文。), Google学术引用2102次, Web of Science学术引用 1500+次。

Broad Learning Systems development history

BLS is a shallow NN, It features efficient training and dynamic incremental growth. It is an efficient solution for lightweight deployment at the edge to achieve collaboration between large and small models.

domestic and international recognition:
Stanford University, University of Oxford,
Tsinghua University etc. 110+ research
institutes, published 1500+ BLS related
papers, more than 300 patents filed.



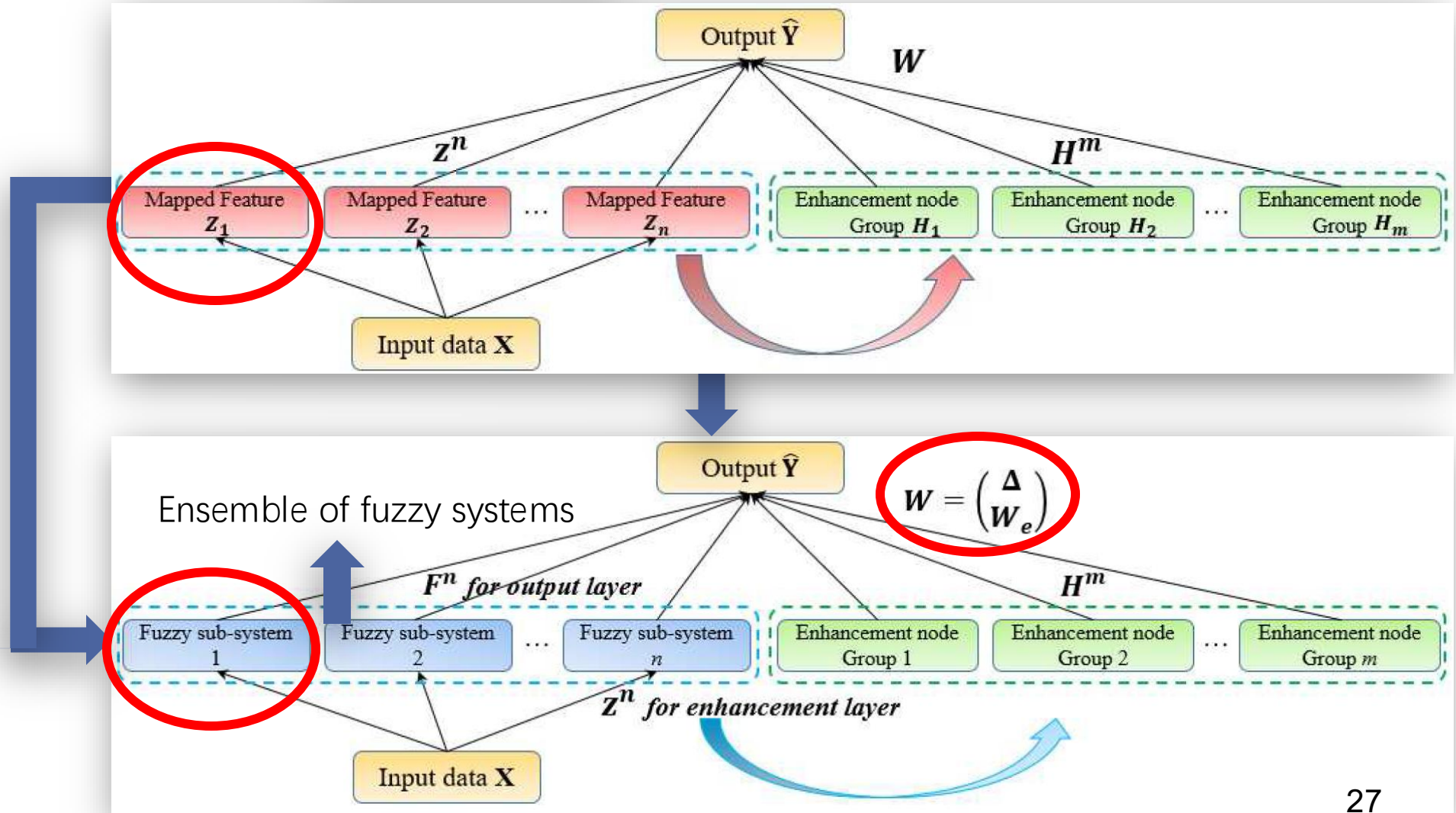
X. Gong, T. Zhang*, C. L. P. Chen and Z. Liu, "Research Review for Broad Learning System: Algorithms, Theory, and Applications," in IEEE Transactions on Cybernetics, doi: 10.1109/TCYB.2021.3061094.

2

Structure of a Fuzzy Broad Learning System

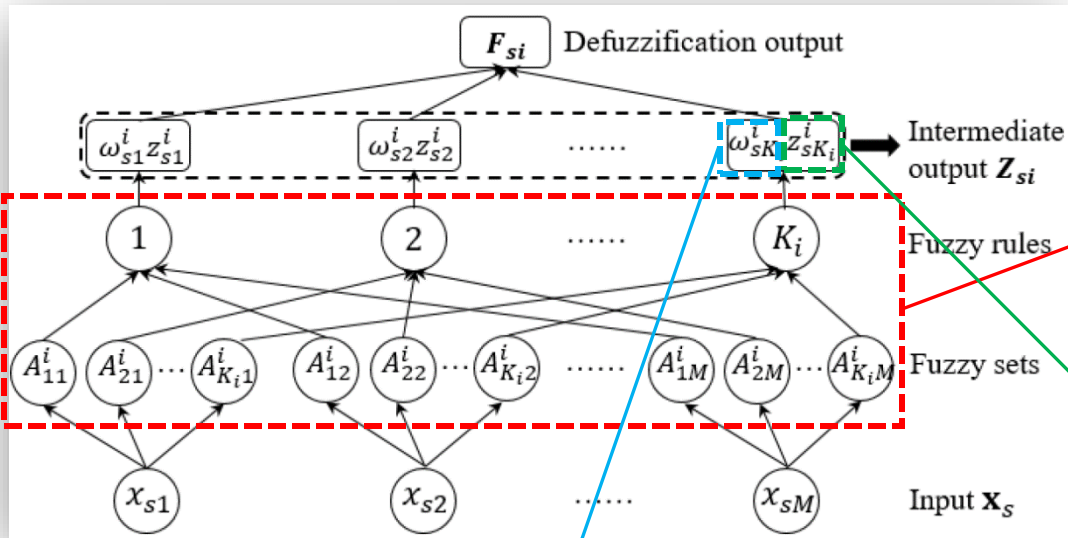
■ Broad Learning System → Fuzzy Broad Learning System

- Remove the sparse-autoencoder
- Replace the feature nodes with Takagi-Sugeno fuzzy sub-systems





Structure of the i-th fuzzy sub-system



Suppose that there are n fuzzy sub-systems and m groups of enhancement nodes in a Fuzzy BLS. The input data are $\mathbf{X} = (x_1, x_2, \dots, x_N)^T \in \mathbb{R}^{N \times M}$ where $x_s = (x_{s1}, x_{s2}, \dots, x_{sM})$, $s = 1, 2, \dots, N$. Suppose that there are K_i fuzzy rules in the i th fuzzy sub-system which have the following form

If x_{s1} is A_{k1}^i and x_{s2} is $A_{k2}^i \dots$ and x_{sM} is A_{kM}^i then $z_{sk}^i = f_k^i(x_{s1}, x_{s2}, \dots, x_{sM})$, $k = 1, 2, \dots, K_i$.

We adopt the first-order TS fuzzy system and let

$$z_{sk}^i = f_k^i(x_{s1}, x_{s2}, \dots, x_{sM}) = \sum_{t=1}^M \alpha_{kt}^i x_{st}, \quad (7)$$

where α_{kt}^i is the coefficient.

The fire strength of the k th fuzzy rule in the i th fuzzy sub-system is

$$\tau_{sk}^i = \prod_{t=1}^M \mu_{kt}^i(x_{st}), \quad (8)$$

and we denote the weighted fire strength for each fuzzy rule as

$$\omega_{sk}^i = \frac{\tau_{sk}^i}{\sum_{k=1}^{K_i} \tau_{sk}^i}. \quad (9)$$

Source codes link:

<http://dsail.vip/bls.html>

相关论文链接

- [Universal approximation capability of broad learning system and its structural variations](#)
- [Research review for broad learning system: Algorithms, theory, and applications](#)
- [Stacked broad learning system: From incremental flatted structure to deep model](#)

Source codes: ➡ 代码

-
- [BLS](#)

Field	Subfield	Method
Computer vision and image processing	Face recognition	Zhang et al. [144], P-BLS [145], ICBL [146] vision-brain [147], Zhang et al. [148], Luan et al. [149]
	Action recognition	DW-Net [136], Fisher-BLS [150], Sun et al. [151] Wavelet BLS [152], Lei et al. [153], Li et al. [154] FBLS-MD [155], Lin et al. [156], Luo et al. [157]
	Image classification	DABL [94], GBLS [93], Guo et al. [158], Cai et al. [159] AdaBoost-BLS [160], DBNet [161], Dang et al. [162] Distributed BLS [163], IWBLs [164], SS-BLS [115], UBL [125] BLSCF [165], Han et al. [166], EAF-BAP [167], Wang et al. [168]
	Image processing	Low-rank BLS [169], Guo et al. [83], DCBLS [107] Jin et al. [170], SID-BLS [84], Shrivastava et al. [171] SAE-BLS [172], Lei et al. [173], Chen et al. [174], Zhai et al. [175]
	Electroencephalography (EEG)	GCB-Net [137], CNBLS [176], GSS-BLS [118], TCBLS [177] MvBLS [87] Mixed-Norm BLS [178], BDGLS [168] LPVG-based BLS [179], CTW-BLS [180], Zou et al. [181]
Medical/biomedical data analysis	Medical data analysis	LPI-BLS [101], Luo et al. [182], MEKM-BLS [183] MISSIM [184], FNT [185], Chew et al. [186], CFCEBLS [187]
	Medical image segmentation	Liu et al. [188], Atlas-BLS [189]
System modeling and prediction	Time series prediction	MIE-BLS [190], R-BLS [64], SM-BLS [95], EMD-BLS [191] BLS-EC [105], Spatio-Temporal BLS [192], Zhou et al. [193] RM-BLS [96], SPCA-BLS [194], Yang et al. [195]
	System modeling	Weight BLS [77], CCBLS [196], Kuok et al. [197] ProBLS [198], Sei-BLS [199], RDBLS [200], BCRBM [201] Regularized-BLS [53], MRBL [202], Lai et al. [203]
Internet and communication engineering	Network and information security	RBF-BLS [204]–[206], Ma et al. [207] Zhong et al. [208], BL-TDA [209], ASSD [210]
	Internet and smart city	Wei et al. [211], BRL [212] Broad Forest [213], Xu et al. [214]
	Robotics	BLS-FDDEIF [215], BLS-SMC [216], BLS-FONTS [217]
Control and robotics	System control	Adaptive-BLS [218], RIFP [219], BFNN [85] DT-BLS [220], Feng et al. [80], Yuan et al. [221] WBLAF [36], [39], QBLS [38], TDFW-BLS [37]
	Motors fault diagnosis	BCNN [74], PABSFD [222], bagging-BLS [104] Wang et al. [223], Jiang et al. [224], transfer-BLS [225] Multi-stage BLS [226], FIBL [227], Han et al. [228], OBLS [229]
Nature language processing	Text classification	Gate-BLS [63], BMT-Net [135], BroXLNet [138]

信息科学领域位居前十位的热点前沿主要集中于人工智能基础理论方法、6G通信、人-机交互、类脑智能、医学信息处理等方向。人工智能基础理论方法方面的热点前沿有5个，生成式对抗网络、宽度学习系统、用于边缘计算的联邦学习成为新的热点前沿，可解释人工智能从去年的新兴前沿成为今年的热点前沿；强化学习相关前沿多次出现在热点前沿中，本期重点是推动强化学习解决真实世界问题的MuZero算法。6G通信方面的热点前沿有两个，深度学习在物理层通信中的应用首次成为热点前沿，可重构智能超表面是从去年的新兴前沿进入到今年的热点前沿。在人-机交互方面，下一代VR/AR实时全息近眼显示方法首次成为热点前沿。在类脑智能方面，脉冲神经网络及其神经形态芯片首次出现。在医学信息处理方面，用于脑电信号分析的卷积神经网络首次成为热点前沿。

表 52 信息科学领域 Top 10 热点前沿

排名	热点前沿	核心论文	被引频次	核心论文平均出版年
1	用于边缘计算的联邦学习	22	3682	2020.2
2	宽度学习系统	6	1053	2020.0
3	可重构智能超表面	32	9372	2019.7
4	下一代 VR/AR 实时全息近眼显示方法	3	457	2019.3
5	可解释人工智能	4	2900	2019.0
6	脉冲神经网络及其神经形态芯片	13	2931	2018.6
7	深度学习在物理层通信中的应用	13	2949	2018.5
8	生成式对抗网络	8	15051	2018.4
9	MuZero 强化学习算法	6	3607	2018.3
10	用于脑电信号分析的卷积神经网络	9	2531	2018.2

喜报

2023年11月28日，中国科学院科技战略咨询研究院、中国科学院文献情报中心与科睿唯安联合发布的《2023研究前沿》报告和《2023研究前沿热度指数》报告，遴选出2023年全球较为活跃或发展迅速的128个研究前沿，并对相关学科的发展趋势和重点问题进行了研判。

中国自动化学会副理事长、华南理工大学教授陈俊龙首创的“宽度学习系统”，荣获2023年信息科学领域热点前沿第二名！

表 52 信息科学领域 Top 10 热点前沿

排名	热点前沿	核心论文	被引频次	核心论文平均出版年
1	用于边缘计算的联邦学习	22	3682	2020.2
2	宽度学习系统	6	1053	2020.0
3	可重构智能超表面	32	9372	2019.7
4	下一代 VR/AR 实时全息近眼显示方法	3	457	2019.3
5	可解释人工智能	4	2900	2019.0
6	脉冲神经网络及其神经形态芯片	13	2931	2018.6
7	深度学习在物理层通信中的应用	13	2949	2018.5
8	生成式对抗网络	8	15051	2018.4
9	MuZero 强化学习算法	6	3607	2018.3
10	用于脑电信号分析的卷积神经网络	9	2531	2018.2

Wu WenJun AI Outstanding Contribution Award, China

首提“宽度学习系统”，实现边缘端智能学习

Dynamic deep neural networks (Based on incremental stacked BLS)

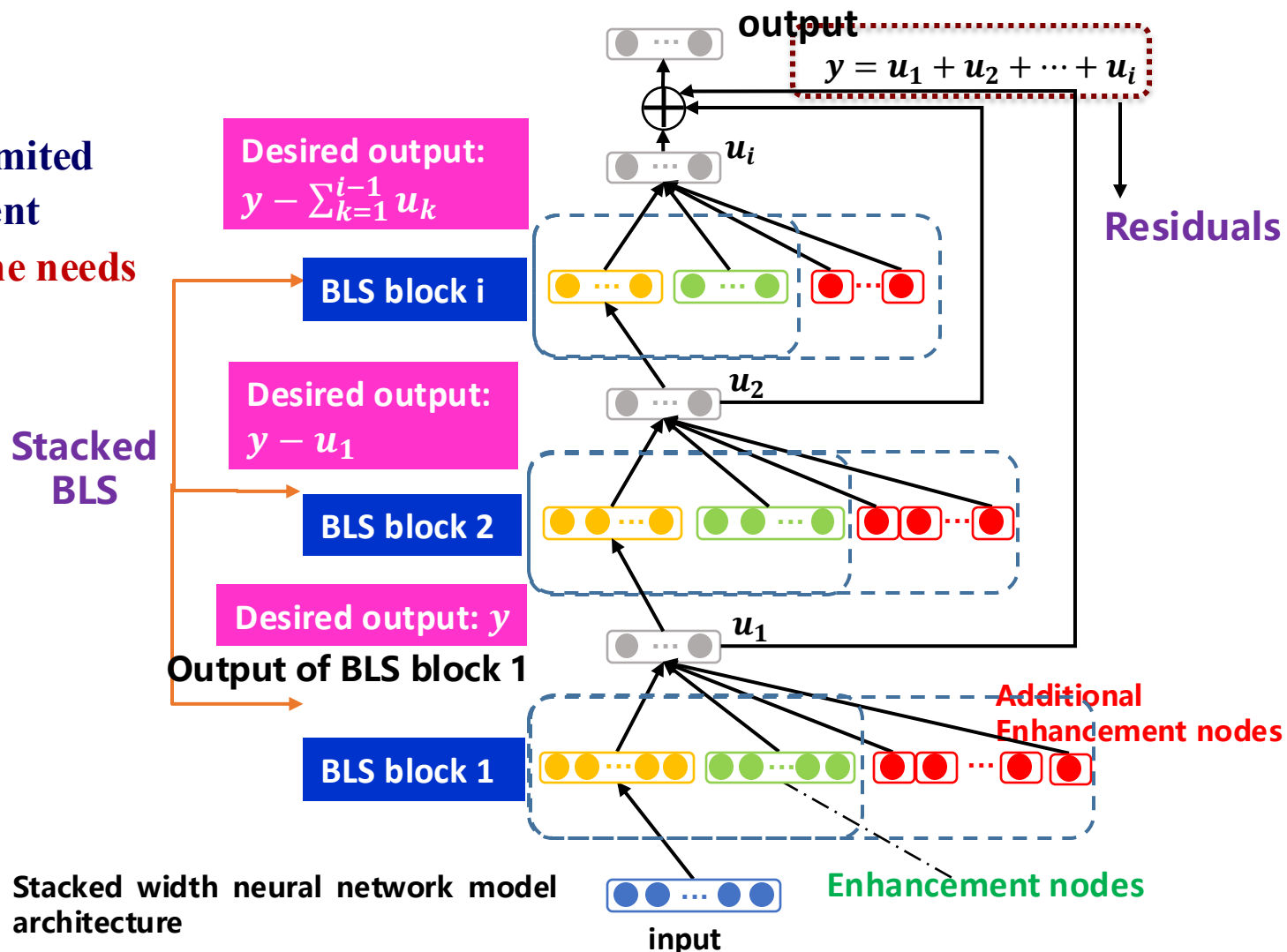
Challenge

Currently, network structures have limited precision, **fixed frameworks**, and insufficient learning capabilities. This does not **meet the needs of brain-inspired research**, and innovative computing frameworks need to be further designed based on brain-related models.

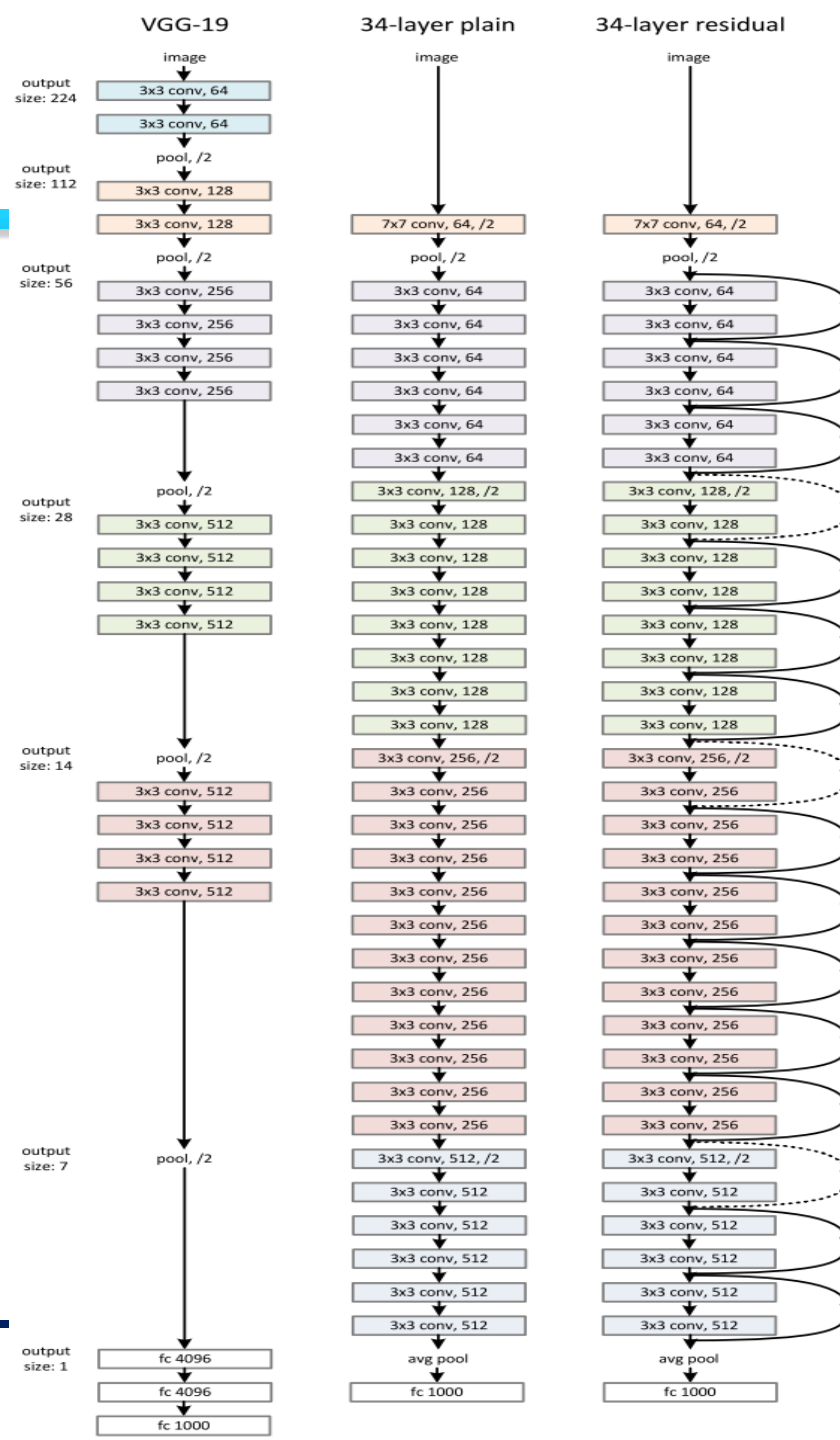
Innovation

Stacking multiple BLS blocks forms a deep network structure.

Introducing the concept of residuals,



ResNet



Compared with ResNet Structure

TABLE XII: Performance comparison between Stacked CCFBLS and ResNet in SVHN

Algorithms	ResNet-50	CCFBLS	Stacked CCFBLS
Testing Accuracy	92.06%	96.69%	97.12%
No. of blocks	-	-	4
No. of layers	50	26	35
No. of Param.	23.521M	2.934M	2.959M



Fig. 12: Example images selected from SVHN Dataset

TABLE XI: Performance comparison between Stacked CCFBLS and ResNet in CIFAR-10

Algorithms	ResNet-32	CCFBLS	Stacked CCFBLS
Testing Accuracy	92.49%	94.29%	94.78%
No. of blocks	-	-	3
No. of layers	32	26	32
No. of Param.	11.174M	2.933M	2.951M

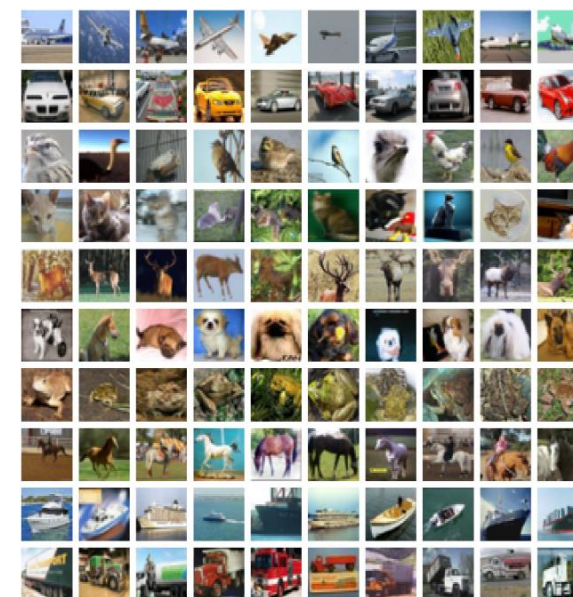


Fig. 11: Example images selected from CIFAR-10 Dataset

Compared with ResNet Structure

TABLE XIII
PERFORMANCE COMPARISON BETWEEN STACKED CCFBLS AND
RESNET IN CIFAR100

Algorithms	ResNet-50	CCFBLS	Stacked CCFBLS
Testing Accuracy	78.51%	81.89%	88.56%
No. of blocks	-	-	3
Accumulative No. of layers	50	28	32
No. of Param.	23.706M	3.658M	3.803M

TABLE XIV
PERFORMANCE COMPARISON WITH SELECTED STATE-OF-THE-ART
DEEP METHODS

Algorithms	CIFAR-10	CIFAR-100	SVHN
Capsule Net [56]	89.3%	-	95.7%
GDCNN [57]	89.23%	66.7%	-
Meta Net [58]	88%	-	-
CGap [59]	93.59%	73%	96.25%
SReluDCNN [60]	93.02%	70.9%	-
Stacked CCFBLS	94.78%	88.56%	97.12%

脑认知的启发：正电子发射断层扫描脑成像 (PET)

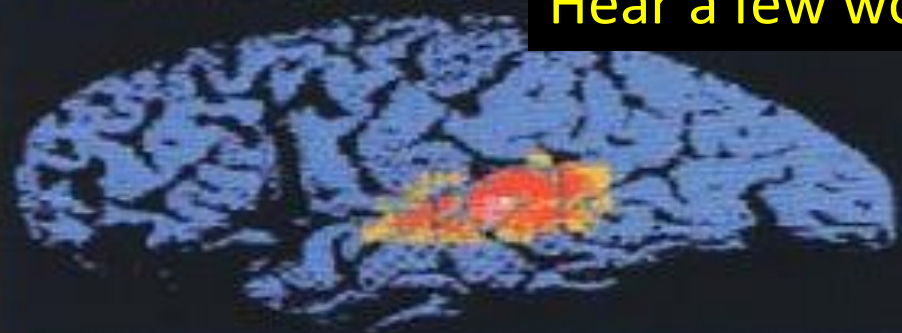
A Looking at words

See a few words



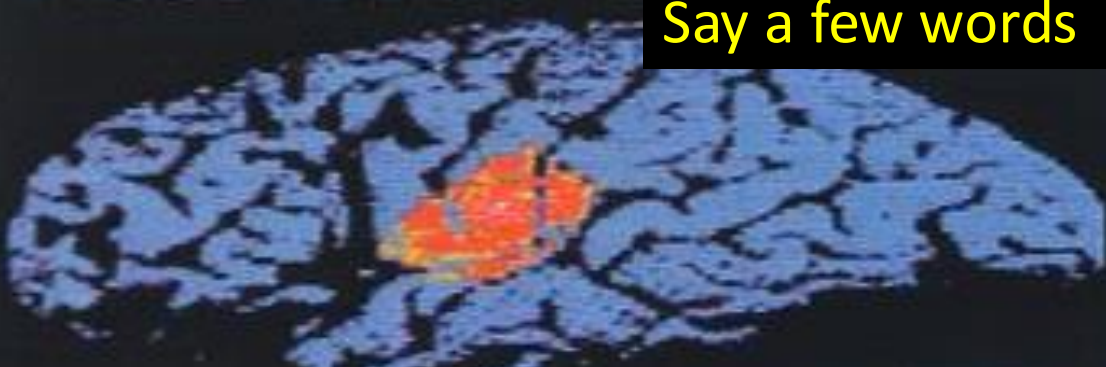
B Listening to words

Hear a few words



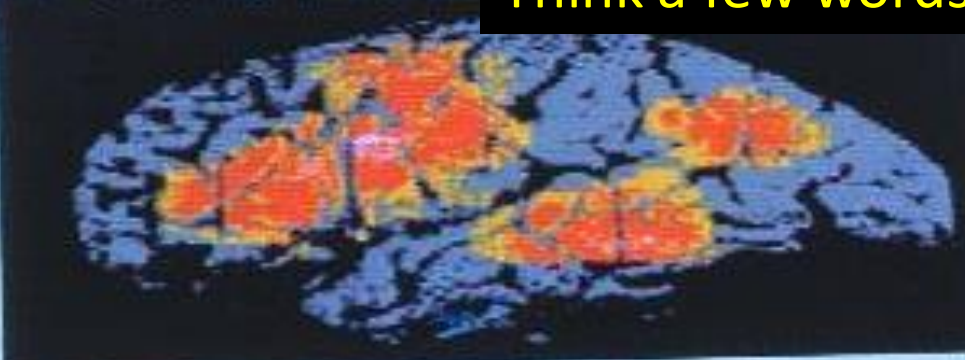
C Speaking words

Say a few words



D Thinking of words

Think a few words



Brain Advance Access published July 14, 2016

doi:10.1093/brain/aww143

BRAIN 2016: Page 1 of 15 | 1

BRAIN
A JOURNAL OF NEUROLOGY

Aside: Robotic-Manipulator Control (advantage of Incremental Learning)

8608

IEEE TRANSACTIONS ON INDUSTRIAL ELECTRONICS, VOL. 67, NO. 10, OCTOBER 2020



3824

Motor Learning and Generalization Using Broad Learning Adaptive Neural Control

Haohui Huang¹, Tong Zhang¹, Member, IEEE, Chenguang Yang¹, Senior Member, IEEE, and C. L. Philip Chen², Fellow, IEEE

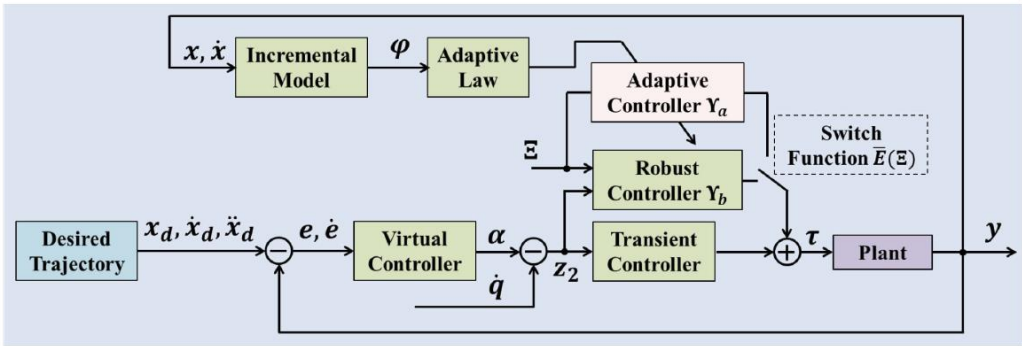


Fig. 2. Control diagram of an adaptive neural controller using a BLS.

IEEE TRANSACTIONS ON CYBERNETICS, VOL. 51, NO. 7, JULY 2021

Optimal Robot–Environment Interaction Under Broad Fuzzy Neural Adaptive Control

Haohui Huang¹, Member, IEEE, Chenguang Yang¹, Senior Member, IEEE, and C. L. Philip Chen², Fellow, IEEE

3826

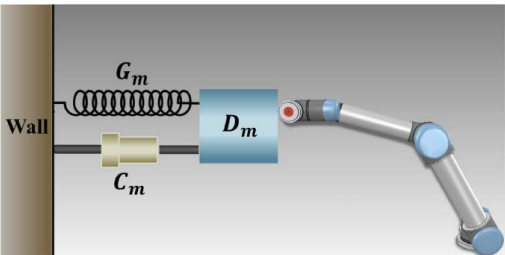


Fig. 1. Impedance model of the robot–environment interaction.

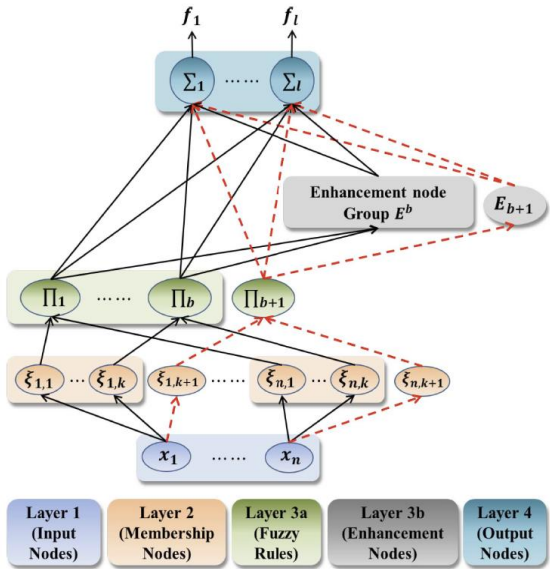
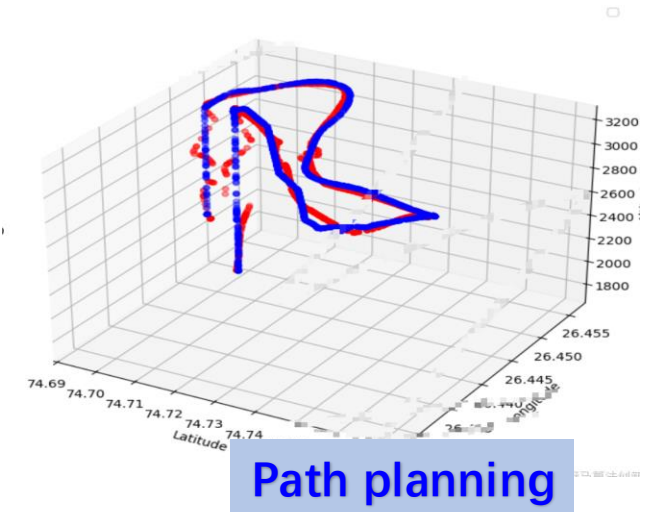
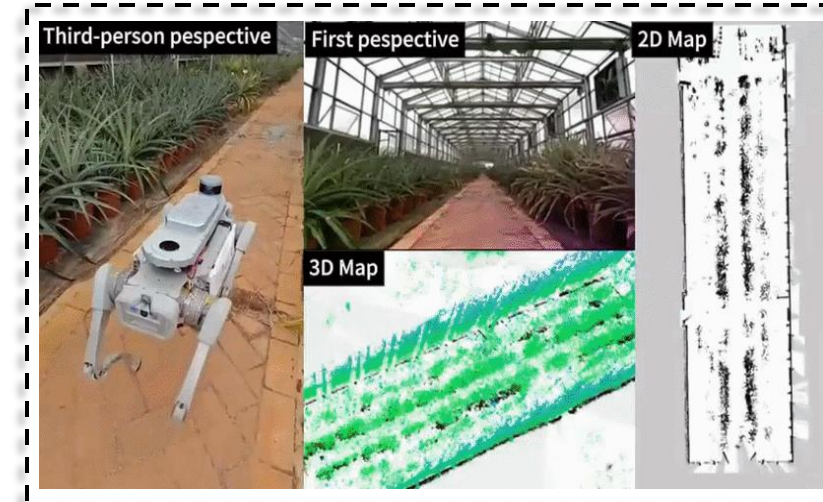


Fig. 4. Schematic of BFNN.



Unmanned System Applications: Incremental Learning in real-time

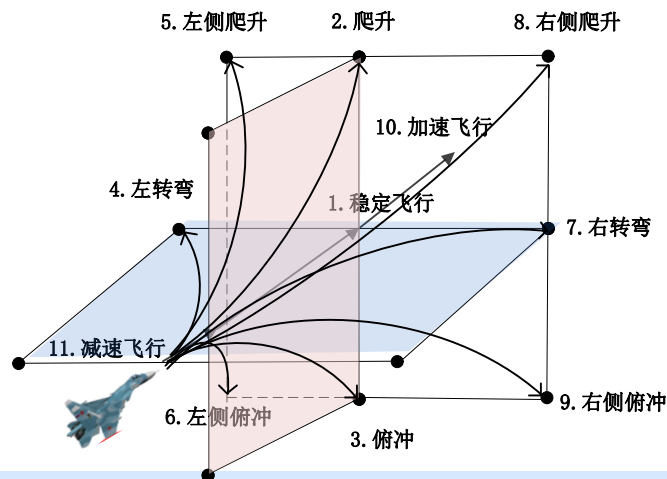


Path planning

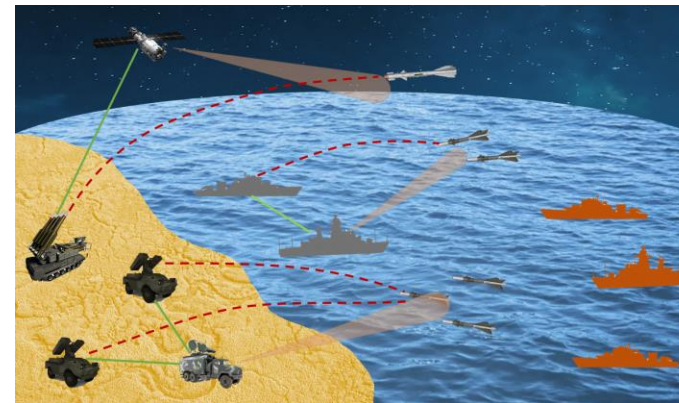
Multi-aircraft adversarial behavior analysis based on a BLS is a key to victory in multi-aircraft combat in complex dynamic environments



Trajectory forecasting



Motion recognition



Intent to reason

Meaning and practical value

- ✓ Anticipate enemy actions in advance
- ✓ Improves attack accuracy
- ✓ Optimize tactical layout
- ✓ Identify false moves and confusing tactics
- ✓ Early warning and avoidance
- ✓ Psychological tactics and psychological advantage
- ✓ Optimize tactical decision-making
- ✓ Auxiliary intelligence and decision support

Inverse-free Broad Learning Systems

Previously, “pseudo-inverse” computation is needed

➤ Least Square Regression

$$\arg \min_W: \|Y - AW\|_2^2 + \lambda \|W\|_2^2$$

$$W = (A^T A + \lambda I)^{-1} A^T Y$$

Greville's Method

All improvements and various transformations are placed on the pseudo-inverse operation of solving W!

$$\begin{cases} B = \begin{cases} C^+, & \text{if } C \neq 0, \\ A^+ D (1 + D^T D)^{-1}, & \text{if } C = 0, \end{cases} \\ C = \tilde{A} - D^T A, \quad D^T = \tilde{A} A^+. \end{cases} \quad \text{其中 } D = (A^m)^+ \xi (Z^n W_{h_{m+1}} + \beta_{h_{m+1}}),$$

Finally, the updated weights $\hat{W}$ are analytically computed as:

$$\hat{W} = W + B(\tilde{Y} - \tilde{A}W).$$

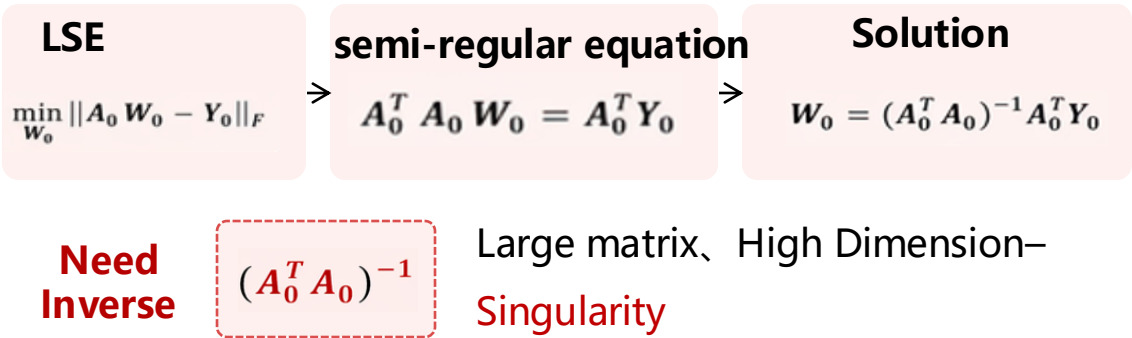
Inverse-free BLS

➤ Least Square Regression

$$\arg \min_W: \|Y - AW\|_2^2 + \lambda \|W\|_2^2$$
$$W = (A^T A + \lambda I)^{-1} A^T Y$$

Problem: inverse operations causes unstable values

Traditional analytical solution solving process



Data Unstable



求解结果对数值敏感、波动大

High computing costs



高维矩阵求逆复杂度高

Real-time capability is limited

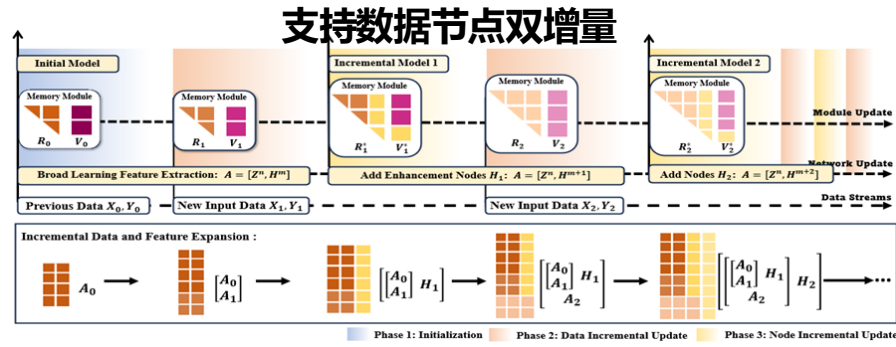
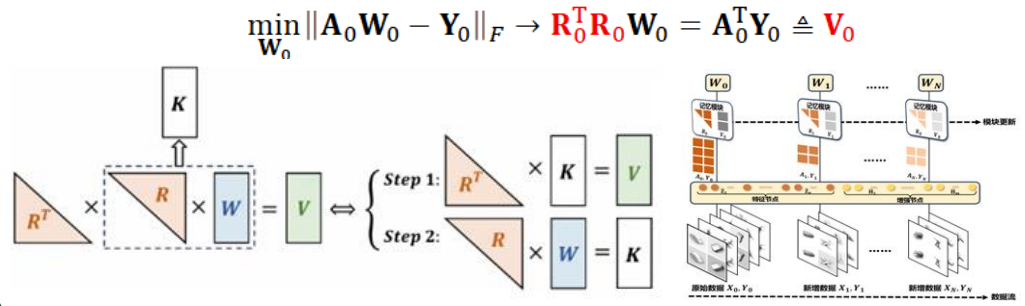


高维矩阵求逆计算时间长

去除求逆运算

Inverse-free improvements: stable and efficient modeling

将最小二乘问题转化为半正规方程求解，借助三角矩阵迭代计算，快速稳定地解析计算更新权重。



对比伪逆方法准确率提升最高2.56%，计算速度提升330.24%！

无逆宽度学习通过消除求逆依赖，实现更稳定、高效的协同建模，更适用于开放复杂环境下云边端智能的实时计算。

Efficient and high-precision incremental calculations: Single round

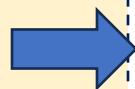
Inverse-free BLS is a breakthrough of data unstable, and achieve, **parallelism** and **lightweighted computing**

A paradigm for calculating data increments without inverted weights

➤ The LSE problem is transformed into a semi-normalization equation

$$\min_{W_0} \|A_0 W_0 - Y_0\|_F \rightarrow$$

To Solve W_0 , multiply A_0^T become $A_0^T A_0$, become $R_0^T R_0$



For large and dense A_0 , previous inverse-based incremental BLS explicitly computes analytical solution W_0 by the pseudo-inverse A_0^+ for initialization:

$$W_0 = A_0^+ Y_0 = (A_0^T A_0 + \lambda I)^{-1} A_0^T Y_0,$$

while storing or updating the pseudo-inverse A_0^+ is inefficient. In contrast, InvF-BLS overcomes these issues via reduced QR decomposition:

$$A_0 = Q_0 R_0,$$

where $Q_0 \in \mathbb{R}^{l_0 \times p}$ is an orthogonal matrix and $R_0 \in \mathbb{R}^{p \times p}$ is an upper triangular matrix.

This transforms the least squares problem into seminormal equations (SNE)[33], which is an easier-solving triangular system:

$$\boxed{A_0^T A_0} W_0 = R_0^T \underbrace{Q_0^T Q_0}_{=I} R_0 W_0 = \boxed{R_0^T R_0} W_0 = A_0^T Y_0 \triangleq V_0. \quad (4)$$

And we set the triangular coefficient matrix and updated labels (R_0, V_0) as the memory module to integrate the new information.

Efficient and high-precision incremental calculations: Single round

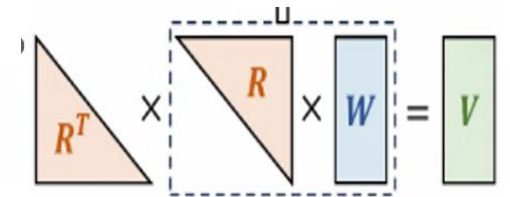
Inverse-free BLS is a breakthrough of data unstable, and achieve, parallelism and lightweighted computing

A paradigm for calculating data increments without inverted weights

- The LSE problem is transformed into a semi-normalization equation

$$\min_{W_0} \|A_0 W_0 - Y_0\|_F \rightarrow R_0^T \underbrace{R_0 W_0}_K = A_0^T Y_0 \triangleq V_0$$

Solve $R_0^T K = V_0$ find K
Then solve $R_0 W_0 = K$ find W_0

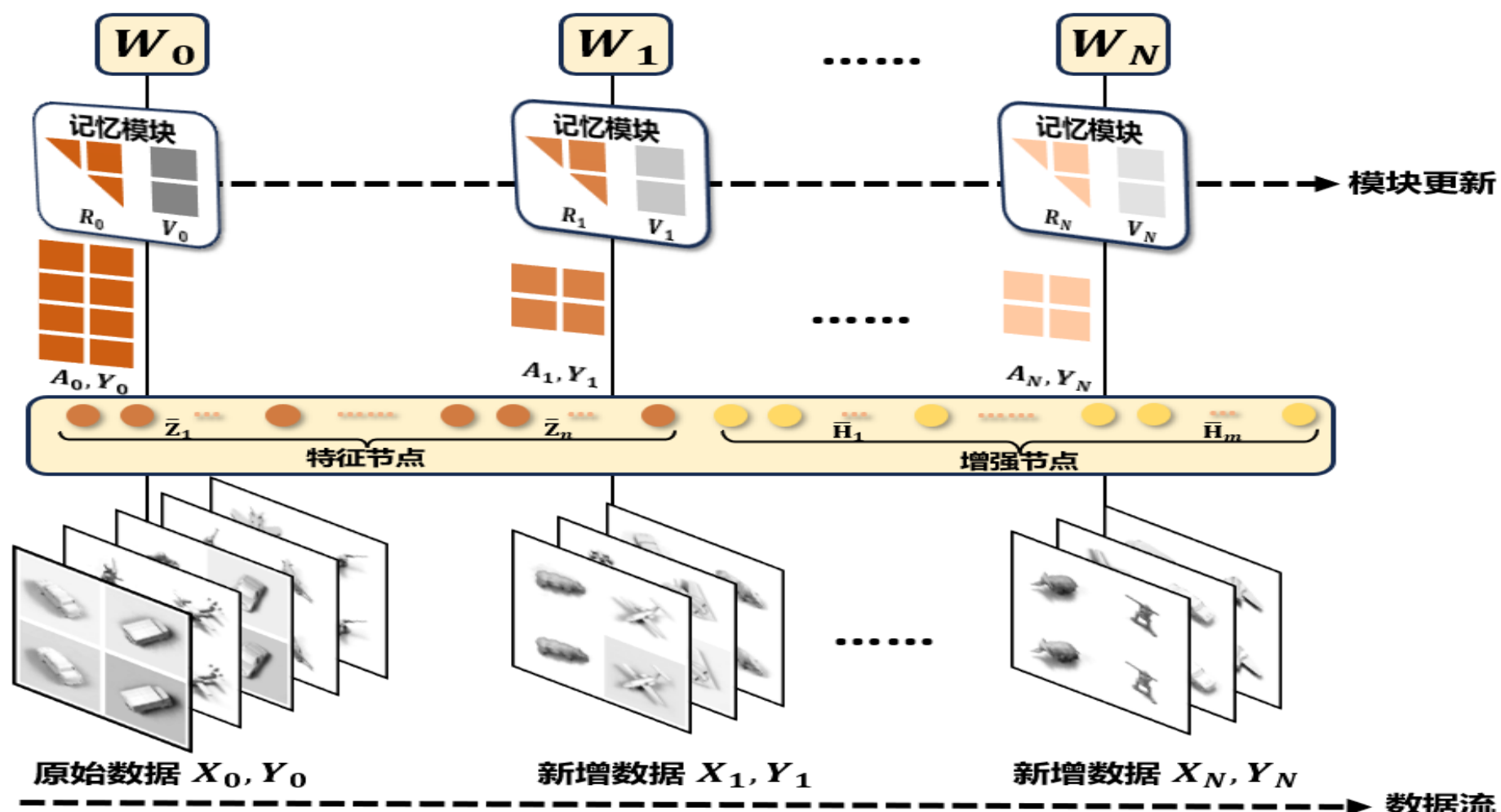


Efficient and high-precision incremental calculations : **Weight updated**

A paradigm for calculating data increments without inverted weights

➤ achieve **parallelism**、**memory module update algorithm** adapted to data increments

➤ By using the **triangular matrix** for forward and backward calculation, the weights are quickly and stably parsed and updated



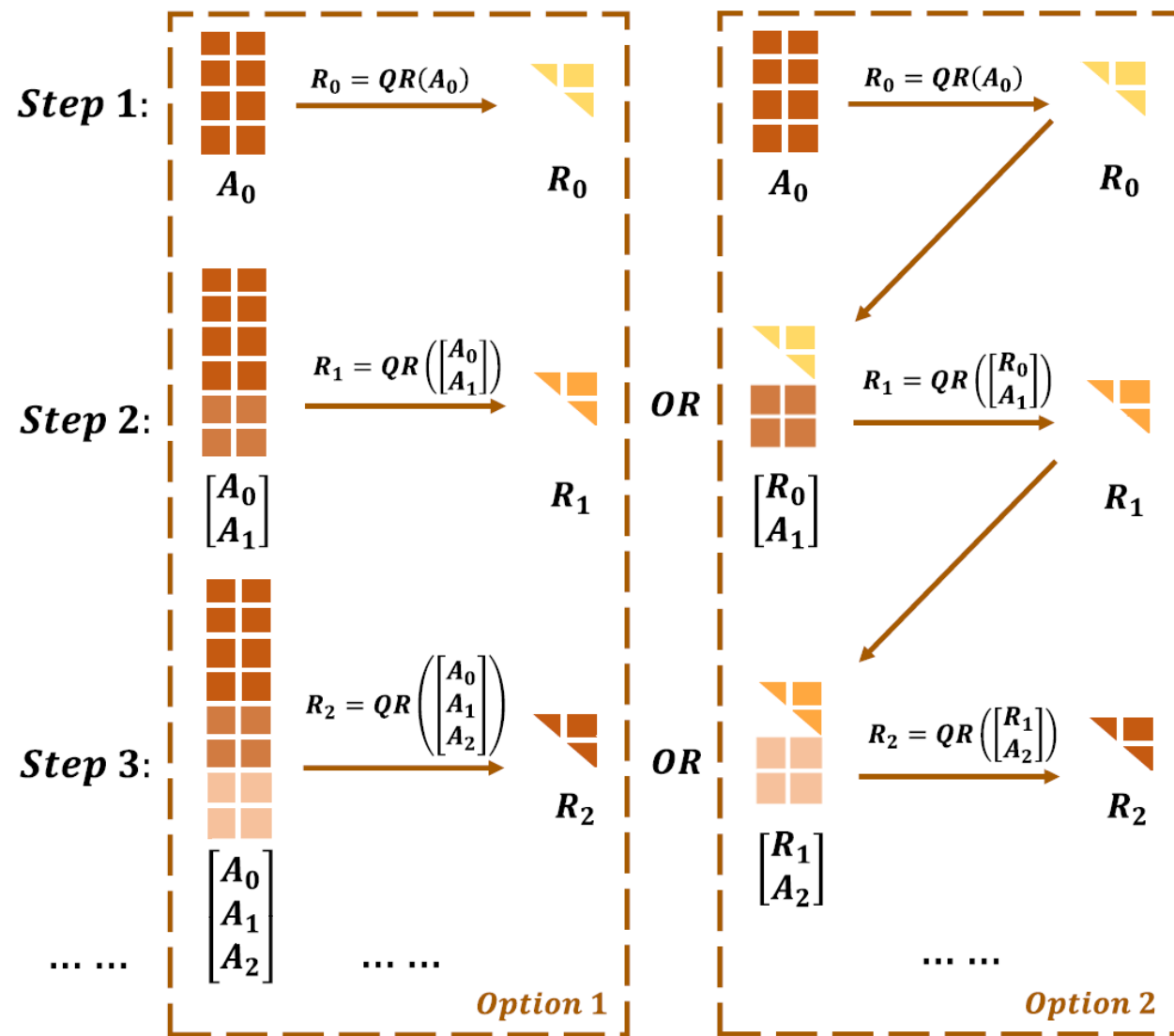


Fig. 4. Incremental R-factor update.

$$\mathbf{W}_k = \mathcal{BS}(\mathbf{R}_k, \mathcal{FS}(\mathbf{R}_k^T, \mathbf{V}_k)),$$

where:

$$\begin{cases} \mathbf{R}_k = \mathcal{QR}\left(\begin{bmatrix} \mathbf{R}_{k-1} \\ \mathbf{A}_k \end{bmatrix}\right), \\ \mathbf{V}_k = \mathbf{R}_{k-1}^T \mathbf{R}_{k-1} \mathbf{W}_{k-1} + \mathbf{A}_k^T \mathbf{Y}_k = \mathbf{V}_{k-1} + \mathbf{A}_k^T \mathbf{Y}_k. \end{cases}$$

This process enables efficient adaptation to recursive update chain:

$$\mathbf{R}_{k-1} \xrightarrow[\text{Data Incremental}]{\text{Update in (7)}} \mathbf{R}_k \quad \text{and} \quad \mathbf{V}_{k-1} \xrightarrow[\text{Data Incremental}]{\text{Update in (8)}} \mathbf{V}_k$$

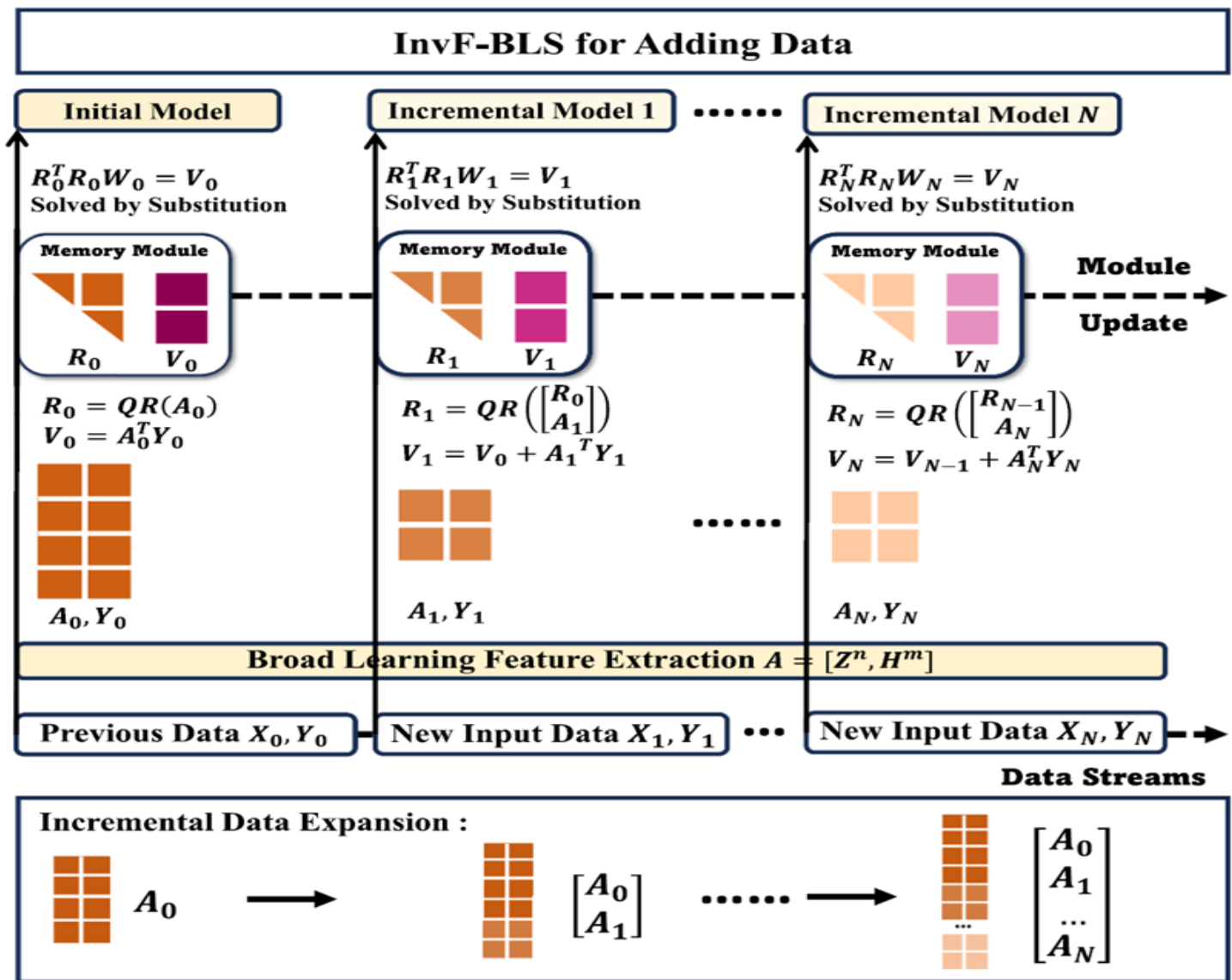
$$\Rightarrow \mathbf{W}_{k-1} \xrightarrow[\text{Update in (9)}]{} \mathbf{W}_k.$$

The overall process of InvF-BLS for adding data is illustrated in Algorithm 3.

Efficient incremental calculations in parallel : Real-time data flow incremental calculation

Data increments

- Through **progressive method**, a **memory matrix update** algorithm is designed to support **data increment** and adapt to **distributed computing** needs under dynamic input.



Summary: Efficient and high-precision data incremental calculation :

$$\mathbf{W}_k = BS(\mathbf{R}_k, FS(\mathbf{R}_k^T, \mathbf{V}_k)),$$

$$\begin{cases} \mathbf{R}_k = QR\left(\begin{bmatrix} \mathbf{R}_{k-1} \\ \mathbf{A}_k \end{bmatrix}\right), \\ \mathbf{V}_k = \mathbf{R}_{k-1}^T \mathbf{R}_{k-1} \mathbf{W}_{k-1} + \mathbf{A}_k^T \mathbf{Y}_k = \mathbf{V}_{k-1} + \mathbf{A}_k^T \mathbf{Y}_k. \end{cases}$$

This process enables efficient adaptation to streaming data by the recursive update chain:

$$\mathbf{R}_{k-1} \xrightarrow[\text{Data Incremental}]{\text{Update in (7)}} \mathbf{R}_k \quad \text{and} \quad \mathbf{V}_{k-1} \xrightarrow[\text{Data Incremental}]{\text{Update in (8)}} \mathbf{V}_k$$

$$\Rightarrow \mathbf{W}_{k-1} \xrightarrow{\text{Update in (9)}} \mathbf{W}_k.$$

Data and nodes increments

- Through **progressive** method, a memory **matrix update** algorithm is designed to support **data increment** and adapt to **distributed computing** needs under **dynamic input**.

2.3.1. Incremental data and feature expansion

InvF-BLS supports flexible expansion sequences (e.g., data-first, node-first, or mixed iterations). The memory module are updated through the expansion order. Fig. 6 gives a data-first-node-second sequence of feature matrix as an example:

$$\begin{aligned} \mathbf{A}_0 &\xrightarrow{\text{Data } \mathbf{X}_1} \begin{bmatrix} \mathbf{A}_0 \\ \mathbf{A}_1 \end{bmatrix} \xrightarrow{\text{Node } \mathbf{H}_1} \begin{bmatrix} \begin{bmatrix} \mathbf{A}_0 \\ \mathbf{A}_1 \end{bmatrix} \mathbf{H}_1 \end{bmatrix} \\ &\xrightarrow{\text{Data } \mathbf{X}_2} \begin{bmatrix} \begin{bmatrix} \mathbf{A}_0 \\ \mathbf{A}_1 \end{bmatrix} \mathbf{H}_1 \\ \mathbf{A}_2 \end{bmatrix} \xrightarrow{\text{Node } \mathbf{H}_2} \begin{bmatrix} \begin{bmatrix} \begin{bmatrix} \mathbf{A}_0 \\ \mathbf{A}_1 \end{bmatrix} \mathbf{H}_1 \\ \mathbf{A}_2 \end{bmatrix} \mathbf{H}_2 \end{bmatrix}. \end{aligned}$$

Scenario 2: Adding Nodes Changes the Factor Size

Adding new columns enlarges the system

When a new node block $\mathbf{H}_q$ is appended, the feature matrix becomes

$$\mathbf{A}_{k+q} = [\mathbf{A}_k \quad \mathbf{H}_q].$$

Accordingly, the regularized Gram matrix changes from

$$\mathbf{M}_k = \mathbf{A}_k^T \mathbf{A}_k + \lambda \mathbf{I}$$

to

$$\mathbf{M}_{k+q} = \mathbf{A}_{k+q}^T \mathbf{A}_{k+q} + \lambda \mathbf{I}.$$

The matrix dimension increases from $p \times p$ to $(p + q) \times (p + q)$. Therefore, the old triangular factor $\mathbf{L}_k$ cannot be used directly for the new system.

$$\mathbf{L}_k \mathbf{L}_k^T = \mathbf{M}_k, \quad \mathbf{L}_{k+q} \mathbf{L}_{k+q}^T = \mathbf{M}_{k+q}.$$

Instead of recomputing a full Cholesky factorization from scratch, we update the factor recursively.

Reuse the old factor

Since $\mathbf{L}_k$ is already available, we keep it in the upper-left block of the new factor and introduce two unknown blocks:

$$\mathbf{L}_{k+q} = \left[\begin{array}{c|c} \mathbf{L}_k & \mathbf{0} \\ \hline \mathbf{E} & \mathbf{G} \end{array} \right].$$

Why this block form?

- $\mathbf{L}_k$ reuses the factorization of the old system.
- $\mathbf{E}$ captures the interaction between the old features and the new nodes.
- $\mathbf{G}$ completes the factorization of the newly introduced part.

Goal: solve $\mathbf{E}, \mathbf{G}$ such that $\mathbf{L}_{k+q} \mathbf{L}_{k+q}^T = \mathbf{M}_{k+q}$.

So the update reduces to solving only two new blocks, rather than factoring the whole matrix again.

Fast Update by Block Matching

$$\mathbf{M}_{k+q} = \left[\begin{array}{c|c} \mathbf{A}_k^T \mathbf{A}_k + \lambda \mathbf{I} & \mathbf{A}_k^T \mathbf{H}_q \\ \hline \mathbf{H}_q^T \mathbf{A}_k & \mathbf{H}_q^T \mathbf{H}_q + \lambda \mathbf{I} \end{array} \right] = \mathbf{L}_{k+q} \mathbf{L}_{k+q}^T = \left[\begin{array}{c|c} \mathbf{L}_k \mathbf{L}_k^T & \mathbf{L}_k \mathbf{E}^T \\ \hline \mathbf{E} \mathbf{L}_k^T & \mathbf{E} \mathbf{E}^T + \mathbf{G} \mathbf{G}^T \end{array} \right].$$

Compare the corresponding blocks

Matching the two matrices gives

$$\mathbf{L}_k \mathbf{E}^T = \mathbf{A}_k^T \mathbf{H}_q \implies \mathbf{E}^T = \mathcal{FS}(\mathbf{L}_k, \mathbf{A}_k^T \mathbf{H}_q),$$

and

$$\mathbf{G} \mathbf{G}^T = (\mathbf{H}_q^T \mathbf{H}_q + \lambda \mathbf{I}) - \mathbf{E} \mathbf{E}^T.$$

The enlarged factor is obtained by

- one triangular solve for $\mathbf{E}$,
- one small Cholesky factorization for $\mathbf{G}$.

Hence, the update is much cheaper than full re-factorization.

Phase 1 (Initialization):

Start by the initiate feature matrix $\mathbf{A}_0$ and labels $\mathbf{Y}_0$, we have $\mathbf{V}_0 = \mathbf{A}_0^T \mathbf{Y}_0$. Setting $\mathbf{L}_0 = \mathbf{R}_0^T$ and applying Cholesky decomposition on

$$\mathbf{M}_0 = \mathbf{A}_0^T \mathbf{A}_0 + \lambda \mathbf{I} = \mathbf{L}_0 \mathbf{L}_0^T \triangleq \mathbf{R}_0^T \mathbf{R}_0.$$

Therefore, $\mathbf{W}_0$ gets from the initial memory module $(\mathbf{R}_0, \mathbf{V}_0)$:

$$\mathbf{L}_0 \mathbf{L}_0^T \mathbf{V}_0 = \mathbf{R}_0^T \mathbf{R}_0 \mathbf{V}_0 = \mathbf{V}_0.$$

Phase 2 (Data Incremental Update):

For new data block $(\mathbf{A}_k, \mathbf{Y}_k)$, the memory module $(\mathbf{R}_k, \mathbf{V}_k)$ is updated as (7) and (8):

$$\begin{cases} \begin{bmatrix} \mathbf{R}_{k-1} \\ \mathbf{A}_k \end{bmatrix} = \mathbf{Q}_k \mathbf{R}_k, \\ \mathbf{V}_k = \mathbf{V}_{k-1} + \mathbf{A}_k^T \mathbf{Y}_k. \end{cases} \quad (10)$$

Moreover, we can conclude

$$\mathbf{L}_k = \mathbf{R}_k^T, \quad 1 \leq k \leq N, \quad (11)$$

where $\mathbf{M}_k = \mathbf{A}_k^T \mathbf{A}_k + \lambda \mathbf{I} = \mathbf{L}_k \mathbf{L}_k^T$, proved in Theorem 3.2 of Section 3. It is a bridge connecting data and node increments.

Phase 3 (Node Incremental Update):

When adding nodes $\mathbf{H}_k$, based on the bright $\mathbf{L}_k = \mathbf{R}_k^T$ and node incremental methods in Chen et al.[34], $\mathbf{L}_k^*$ could be updated by :

$$\mathbf{L}_k^* = \begin{bmatrix} \mathbf{L}_k & 0 \\ \mathbf{E} & \mathbf{G} \end{bmatrix} = \begin{bmatrix} \mathbf{R}_k^T & 0 \\ \mathbf{E} & \mathbf{G} \end{bmatrix} \triangleq (\mathbf{R}_k^*)^T, \quad (12)$$

where

$$\begin{cases} \mathbf{E}^T = \mathcal{BS}(\mathbf{L}_k, \mathbf{A}_{0:k}^T \mathbf{H}_k), \\ \mathbf{G} \mathbf{G}^T = \mathbf{H}_k^T \mathbf{H}_k - \mathbf{E} \mathbf{E}^T + \lambda \mathbf{I}. \end{cases}$$

And $\mathbf{V}_k^*$ can be computed from $\mathbf{V}_k$:

$$\mathbf{V}_k^* = [\mathbf{A}_{0:k}, \mathbf{H}_k]^T \mathbf{Y}_{0:k} = \begin{bmatrix} \mathbf{V}_k \\ \mathbf{H}_k^T \mathbf{Y}_{0:k} \end{bmatrix}. \quad (13)$$

Summary: Efficient and high-precision nodes increments (Accuracy approximation)

$$\mathbf{W}_k = BS\left(\left(\mathbf{L}_k^*\right)^T, \mathcal{FS}\left(\mathbf{L}_k^*, \mathbf{V}_k^*\right)\right),$$

where:

$$\begin{cases} \mathbf{L}_k^* = \begin{bmatrix} QR\left(\begin{bmatrix} \mathbf{R}_{k-1} \\ \mathbf{A}_k \end{bmatrix}\right)^T & 0 \\ \mathbf{E} & \mathbf{G} \end{bmatrix}, \\ \mathbf{V}_k^* = \begin{bmatrix} \mathbf{R}_{k-1}^T \mathbf{R}_{k-1} \mathbf{W}_{k-1} + \mathbf{A}_k^T \mathbf{Y}_k \\ \mathbf{H}_k^T \mathbf{Y}_{0:k} \end{bmatrix} = \begin{bmatrix} \mathbf{V}_{k-1} + \mathbf{A}_k^T \mathbf{Y}_k \\ \mathbf{H}_k^T \mathbf{Y}_{0:k} \end{bmatrix}. \end{cases}$$

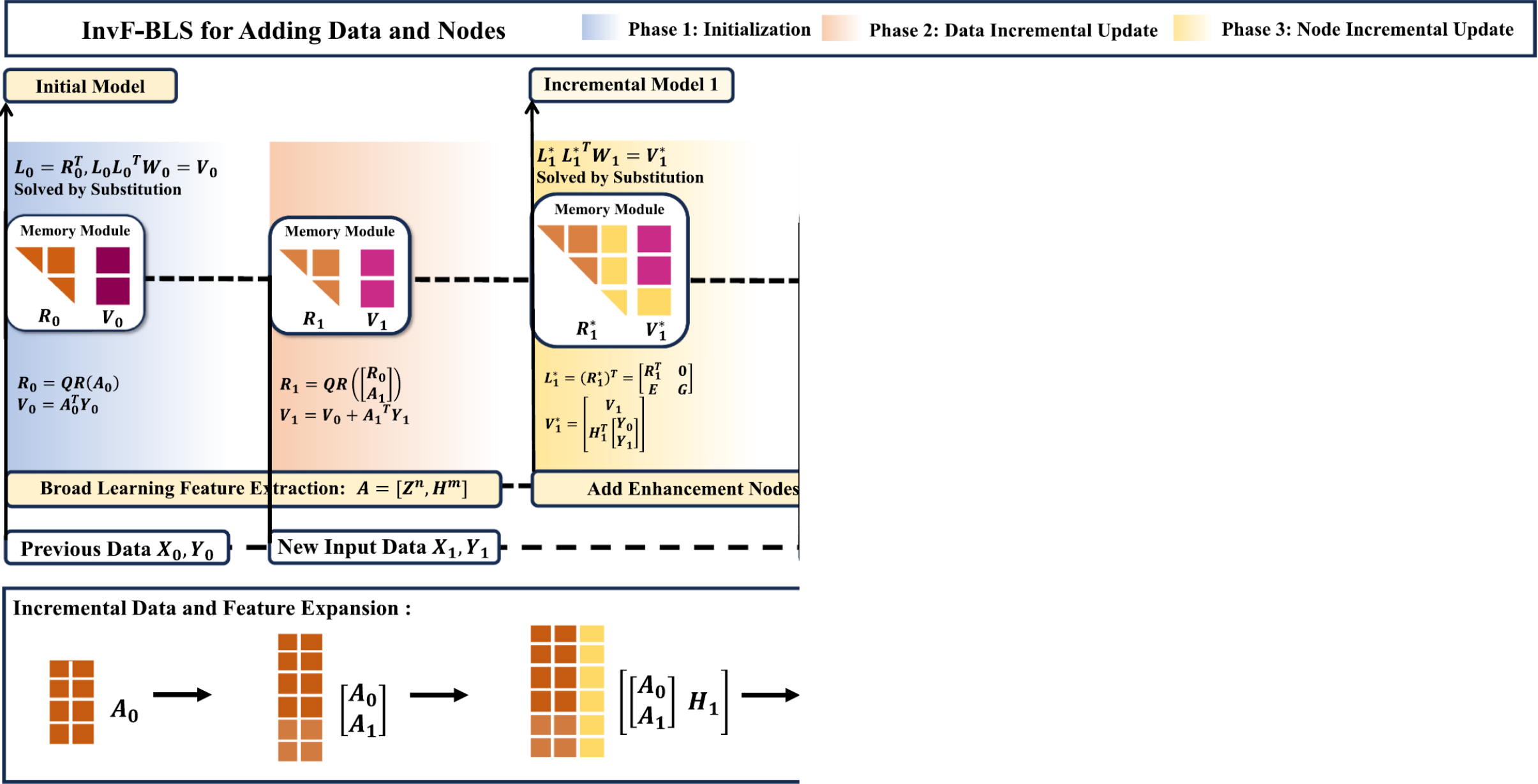
The recursive update process follows a chain:

$$\mathbf{R}_{k-1} \xrightarrow[\text{Data Incremental}]{\text{Update in (10)}} \mathbf{R}_k \xrightarrow[\text{Bridge}]{\text{Transposition in (11)}} \mathbf{L}_k \xrightarrow[\text{Node Incremental}]{\text{Update in (12)}} \mathbf{L}_k^*$$

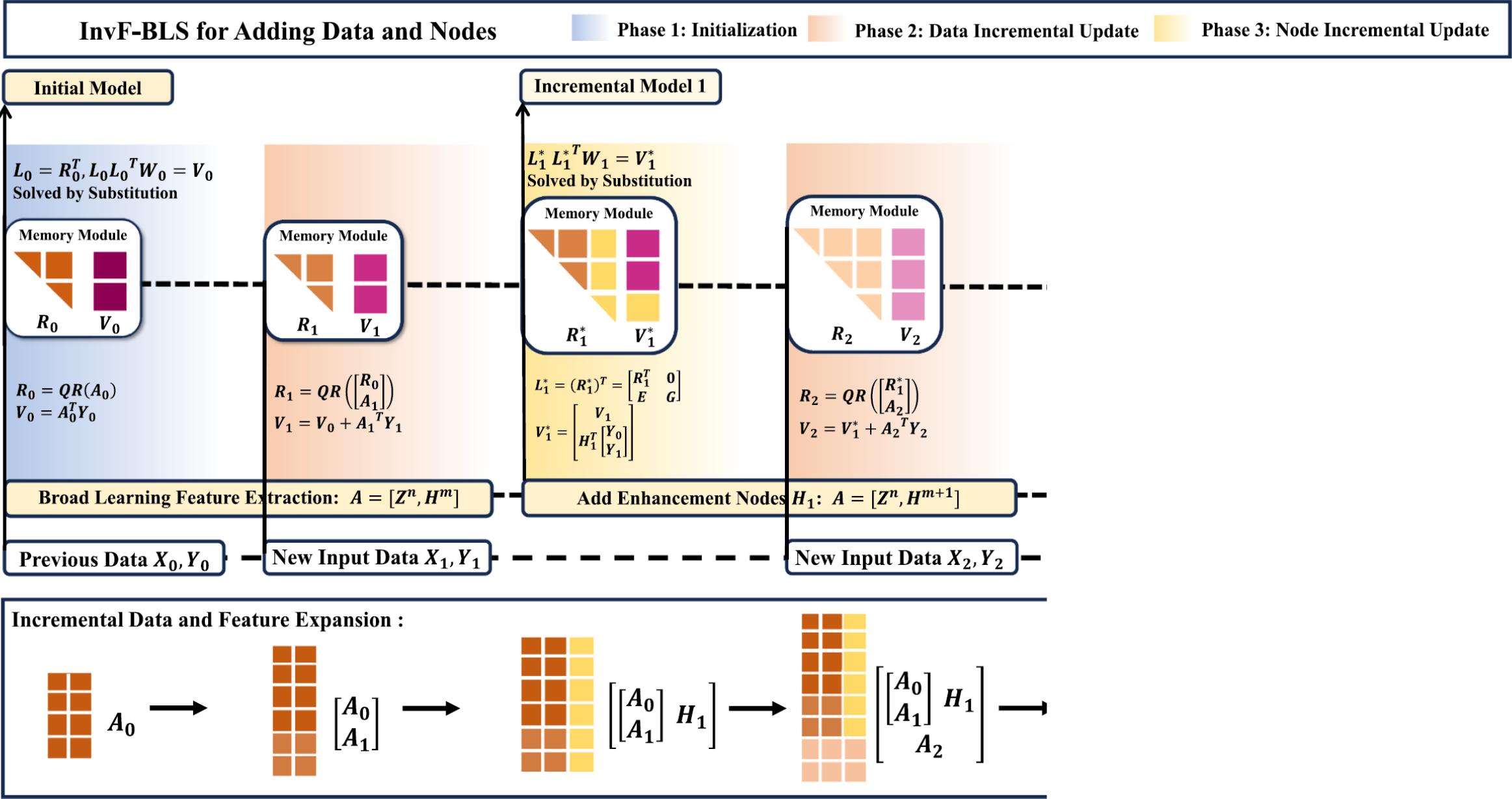
$$\text{and } \mathbf{V}_{k-1} \xrightarrow[\text{Data Incremental}]{\text{Update in (10)}} \mathbf{V}_k \xrightarrow[\text{Node Incremental}]{\text{Update in (13)}} \mathbf{V}_k^*$$

$$\Rightarrow \mathbf{W}_{k-1} \xrightarrow{\text{Update in (14)}} \mathbf{W}_k.$$

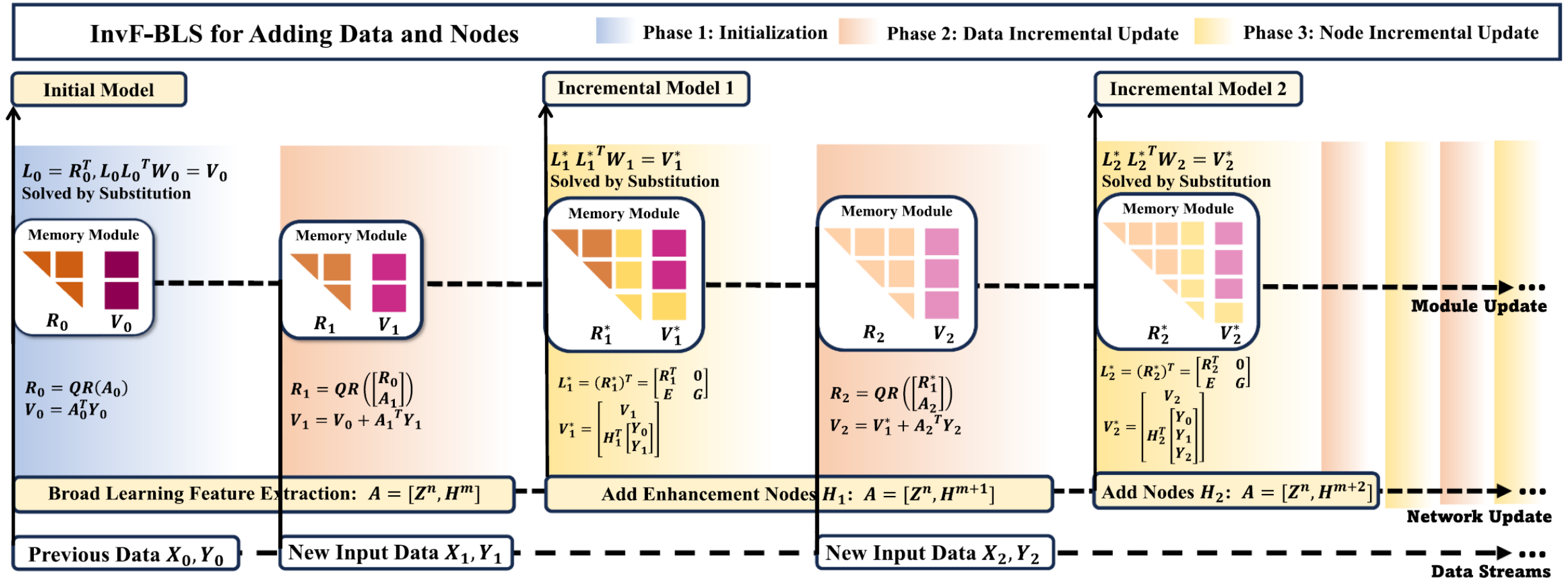
Real-time data flow+nodes incremental calculation



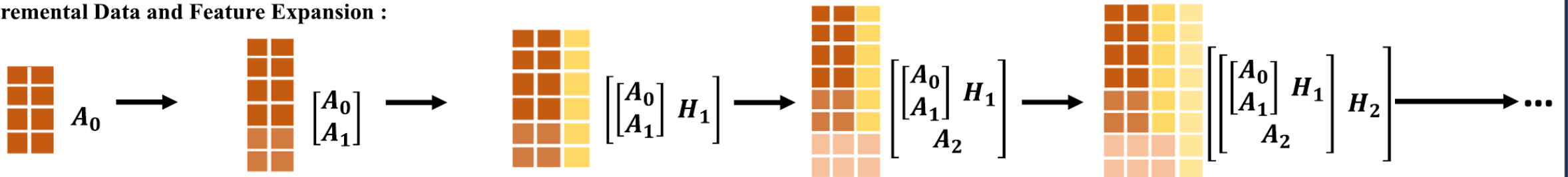
Real-time data flow+nodes+data incremental calculation



Real-time data flow+nodes+data+nodes incremental



Incremental Data and Feature Expansion :



Efficient and high-precision **data** **incremental** calculation

$$\mathbf{W}_k = \mathcal{BS}(\mathbf{R}_k, \mathcal{FS}(\mathbf{R}_k^T, \mathbf{V}_k)),$$

$$\begin{cases} \mathbf{R}_k = \mathcal{QR}\left(\begin{bmatrix} \mathbf{R}_{k-1} \\ \mathbf{A}_k \end{bmatrix}\right), \\ \mathbf{V}_k = \mathbf{R}_{k-1}^T \mathbf{R}_{k-1} \mathbf{W}_{k-1} + \mathbf{A}_k^T \mathbf{Y}_k = \mathbf{V}_{k-1} + \mathbf{A}_k^T \mathbf{Y}_k. \end{cases}$$

This process enables efficient adaptation to streaming data recursive update chain:

$$\mathbf{R}_{k-1} \xrightarrow[\text{Data Incremental}]{\text{Update in (7)}} \mathbf{R}_k \quad \text{and} \quad \mathbf{V}_{k-1} \xrightarrow[\text{Data Incremental}]{\text{Update in(8)}} \mathbf{V}_k$$

$$\Rightarrow \mathbf{W}_{k-1} \xrightarrow{\text{Update in(9)}} \mathbf{W}_k.$$

Efficient and high-precision **node** **increments** (precision approximation)

$$\mathbf{W}_k = \mathcal{BS}\left((\mathbf{L}_k^*)^T, \mathcal{FS}(\mathbf{L}_k^*, \mathbf{V}_k^*)\right),$$

where:

$$\begin{cases} \mathbf{L}_k^* = \begin{bmatrix} \mathcal{QR}\left(\begin{bmatrix} \mathbf{R}_{k-1} \\ \mathbf{A}_k \end{bmatrix}\right)^T & 0 \\ \mathbf{E} & \mathbf{G} \end{bmatrix}, \\ \mathbf{V}_k^* = \begin{bmatrix} \mathbf{R}_{k-1}^T \mathbf{R}_{k-1} \mathbf{W}_{k-1} + \mathbf{A}_k^T \mathbf{Y}_k \\ \mathbf{H}_k^T \mathbf{Y}_{0:k} \end{bmatrix} = \begin{bmatrix} \mathbf{V}_{k-1} + \mathbf{A}_k^T \mathbf{Y}_k \\ \mathbf{H}_k^T \mathbf{Y}_{0:k} \end{bmatrix}. \end{cases}$$

The recursive update process follows a chain:

$$\mathbf{R}_{k-1} \xrightarrow[\text{Data Incremental}]{\text{Update in (10)}} \mathbf{R}_k \xrightarrow[\text{Bridge}]{\text{Transposition in (11)}} \mathbf{L}_k \xrightarrow[\text{Node Incremental}]{\text{Update in (12)}} \mathbf{L}_k^*$$

$$\text{and } \mathbf{V}_{k-1} \xrightarrow[\text{Data Incremental}]{\text{Update in (10)}} \mathbf{V}_k \xrightarrow[\text{Node Incremental}]{\text{Update in (13)}} \mathbf{V}_k^*$$

$$\Rightarrow \mathbf{W}_{k-1} \xrightarrow{\text{Update in (14)}} \mathbf{W}_k.$$

Time Efficiency and Accuracy

Table 5

Results of final testing accuracy and total train time on real-world datasets.

Dataset	Incremental BLS		TiBLS		Approximation Method		RI-BLS		InvF-BLS	
	Acc (%)	Time (s)	Acc (%)	Time (s)	Acc (%)	Time (s)	Acc (%)	Time (s)	Acc (%)	Time (s)
LED	74.05	1838.14	74.05	108.21	74.03	61.93	74.04	28.28	74.08	4.76
CIFAR10	88.59	224.92	88.60	46.98	24.42	24.82	88.59	14.67	88.64	1.71
CIFAR100	64.92	87.58	64.75	18.60	52.46	8.86	64.92	6.25	65.28	0.76
Waveform	82.50	0.31	82.63	0.59	81.25	0.18	82.50	0.17	82.80	0.01
Letter	92.72	8.40	93.40	6.83	91.22	3.66	92.72	2.89	95.30	0.26
Pendigits	98.03	1.01	97.76	1.27	98.02	0.47	98.02	0.38	98.90	0.02
Shuttle	98.61	86.57	98.52	24.07	98.61	6.97	98.61	5.56	99.03	0.64

Time Efficiency and Accuracy

Table 8

Accuracy and training time with different numbers of added enhancement nodes on MNIST.

M_1, M_2	N_1, N_2, N_3	One-shot Accuracy(%)	One-shot Training time(s)	Incremental BLS Accuracy(%)	Incremental BLS Training time(s)	InvF-BLS Accuracy(%)	InvF-BLS Training time(s)
10000, 100	10, 10, 2000	97.73	6.963	95.86	3.469	97.82	2.083
10000, 500	10, 10, 2000	97.90	17.340	96.62	8.586	98.32	5.386
10000, 1000	10, 10, 2000	97.99	37.407	97.04	19.656	98.69	13.016
10000, 1500	10, 10, 2000	97.93	66.411	97.17	34.420	98.77	24.501
10000, 2000	10, 10, 2000	98.28	104.472	97.36	53.863	98.89	40.064
10000, 100	20, 10, 2000	97.83	7.754	96.12	3.735	97.87	2.206
10000, 500	20, 10, 2000	97.93	18.509	96.98	9.054	98.49	5.655
10000, 1000	20, 10, 2000	98.30	39.075	97.28	20.192	98.70	13.616
10000, 1500	20, 10, 2000	98.45	68.173	97.38	35.242	98.91	24.667
10000, 2000	20, 10, 2000	98.07	106.888	97.50	54.417	98.89	40.710

Space and Memory Efficiency

Table 2

Comparison of Number of Operations, Time Complexity, Space Requirements and Space Complexity. Note: N^* is the number of new samples; N is the number of original samples; p is the total number of nodes; c is the number of classes.

Method	Number of Operations (Flops)	Time Complexity	Space Requirements	Space Complexity
InvF-BLS (Ours)	Steps: - Incremental R-Factor Update Eq. (7): N^*p^2 - Right-Hand Side Update Eq. (8): $2N^*pc$ - Weight Calculation Eq. (9): $2p^2c$ Total: $N^*p^2 + 2N^*pc + 2p^2c$	$O(N^*p^2 + N^*pc + p^2c)$	$\frac{1}{2}p^2 + \frac{1}{2}p + 2pc$	$O(p^2 + pc)$
Memory and data volume not depend on N, , It only relates to neuron nodes, p				
RI-BLS [22]	Steps (from [22]): - Memory Matrix U^* Update (Eq. 11): $2N^*p^2$ - Memory Matrix V^* Update (Eq. 12): $2N^*pc$ - Weight Calculation (Eq. 15): $p^3 + 2p^2c$ Total: $p^3 + 2N^*p^2 + 2N^*pc + 2p^2c$	$O(p^3 + N^*p^2 + N^*pc + p^2c)$	$3p^2 + 3pc$	$O(p^2 + pc)$
Lots of N in other approaches				
Approximation Method [46]	$p^3 + 4N^*p^2 + 2Npc + 2N^*pc$	$O(p^3 + N^*p^2 + Npc + N^*pc)$	$2p^2 + 4N^*p + 2Np + N^*p$ $+Np + Nc + N^*c + pc$	$O(p^2 + Np + N^*p + Nc + N^*c + pc)$
Incremental BLS [14]	$p^3 + 6NN^*p + 4N^*p^2 + 4N^*pc$ or $N^*3 + 8NN^*p + 2pN^*2 + 4N^*pc$	$O(p^3 + NN^*p + N^*p^2 + N^*pc)$ or $O(N^*3 + NN^*p + pN^*2 + N^*pc)$	$3N^*p + 2Np + NN^* + N^*c + pc$ or $2N^*2 + 5Np + NN^* + N^*c + pc$	$O(N^*p + Np + NN^* + N^*c + pc)$ or $O(N^*2 + N^*p + Np + NN^* + N^*c + pc)$

Real-time computing of edge embodied intelligence



Offline intelligence

via:

BLS provides

Real-time computing in
a dynamic open environment

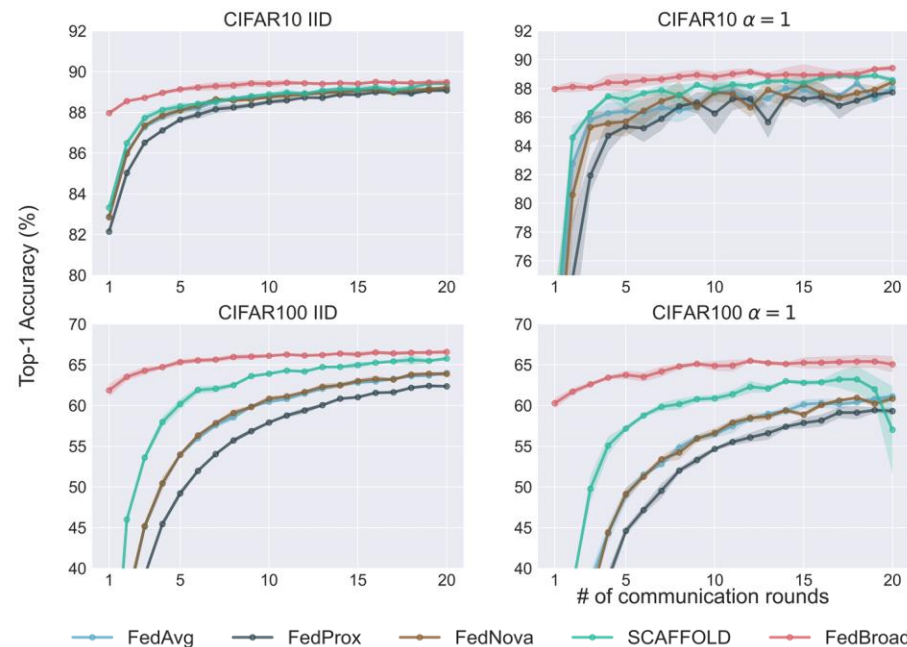
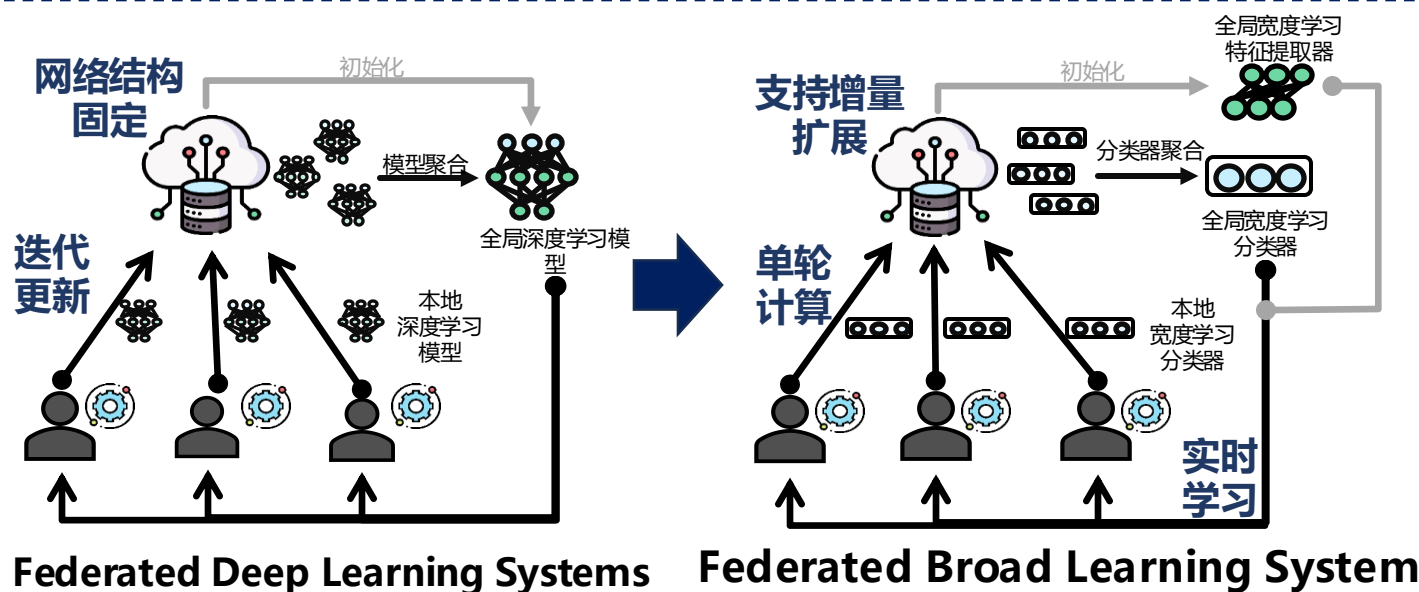


Broad federated learning – a federated modeling approach aimed at **low communication overhead**

effectively solving the problem of excessive communication overhead in traditional federated learning methods and providing a new paradigm for secure distributed deployment of lightweight neural networks in open environments.

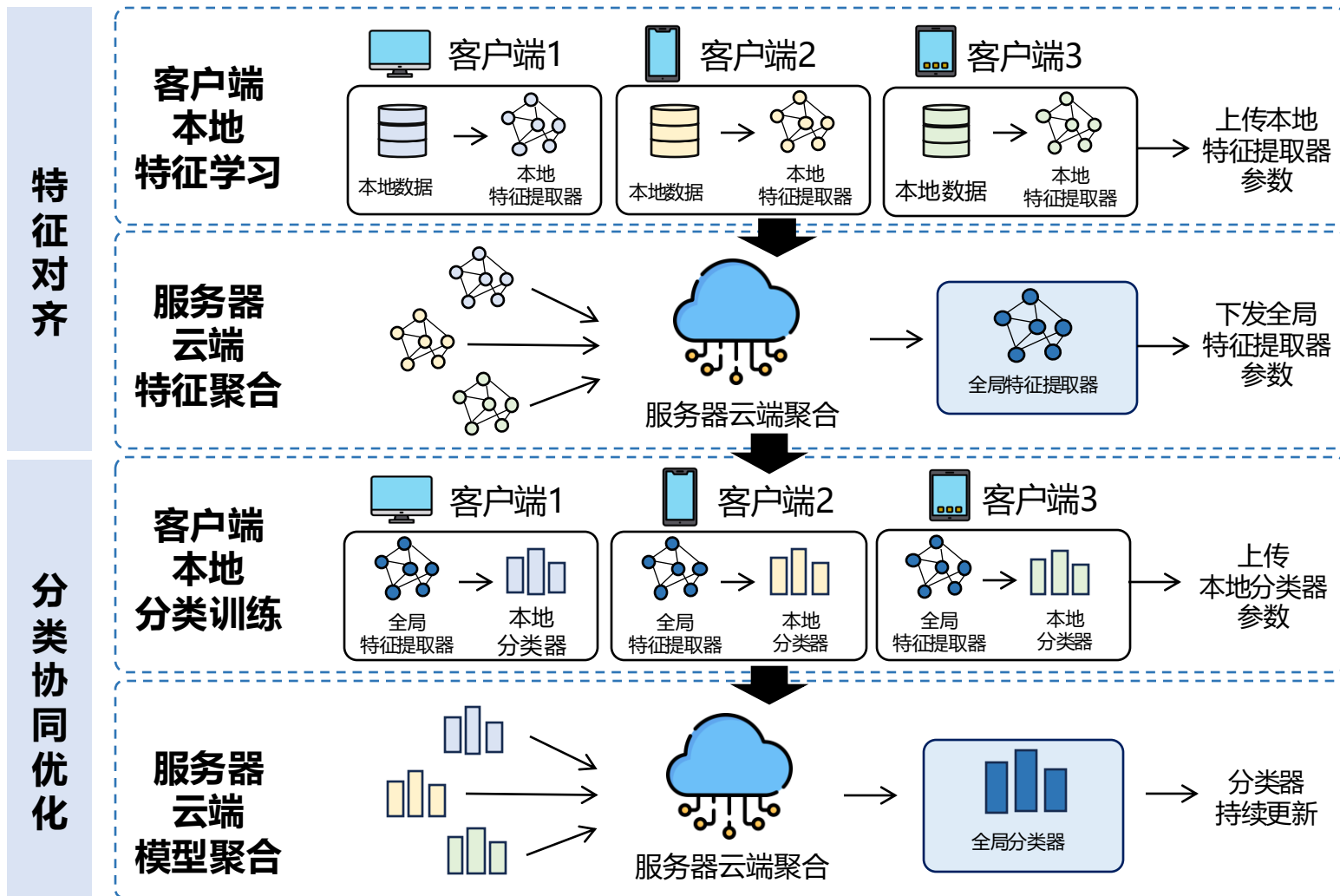
A broad federated learning system that communicates efficiently

- ✓ Leveraging the advantages BLS, feature alignment and near-end regex mechanisms are introduced to enhance the generalization capability of federated systems under heterogeneous data, achieving efficient federated learning for communication.



在保证模型精度的前提下**显著减少通信轮数**

Broad federated learning – a federated modeling approach aimed at **low communication overhead**



Core challenges of traditional federated learning

In an open and complex environment, cloud-edge-end intelligence requires a **privacy-protection collaborative** modeling method that offers lower latency, higher efficiency, and sustainable updates.

Broad federal learning

By utilizing the **single-round analytical training** feature of breadth learning, it replaces traditional federated learning with multi-turn iterative optimization to achieve efficient collaborative modeling.



Data privacy protection
Data never leaves the local area



Reduce multiple rounds of iteration



Low communication overhead
Only transmit model parameters, averaging fewer communication turns **22.9×**倍

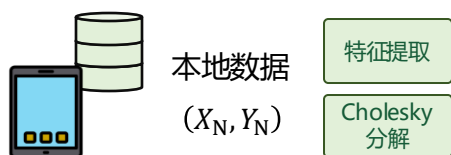
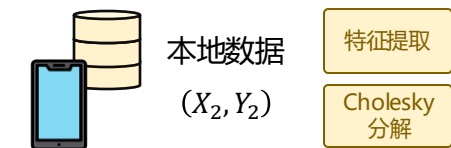
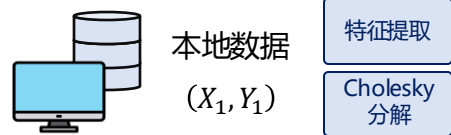


Supports continuous learning
Adapt to a dynamically changing environment

Invert-free Broad Federated Learning

Methodological framework

① Client-side local processing



②

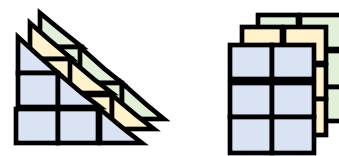
Single round upload

$$A^T Y$$
$$L'$$

$$A^T Y$$
$$L'$$

$$A^T Y$$
$$L'$$

③ Server solves without inverse



Aggregation

$$M = (A_1^T A_1 + \dots + A_n^T A_n)$$
$$W = M^{-1} A_1^T Y_1 + \dots + M^{-1} A_n^T Y_n$$

Global model

Core advantages



Single-round collaborative communication

一次交互完成建模，提升实时协同能力。



Heterogeneous, Distributed, Robust

结果相对数据分布不变，适应复杂场景。



Efficient analytical solutions

避免显式求逆，降低边缘端资源占用。



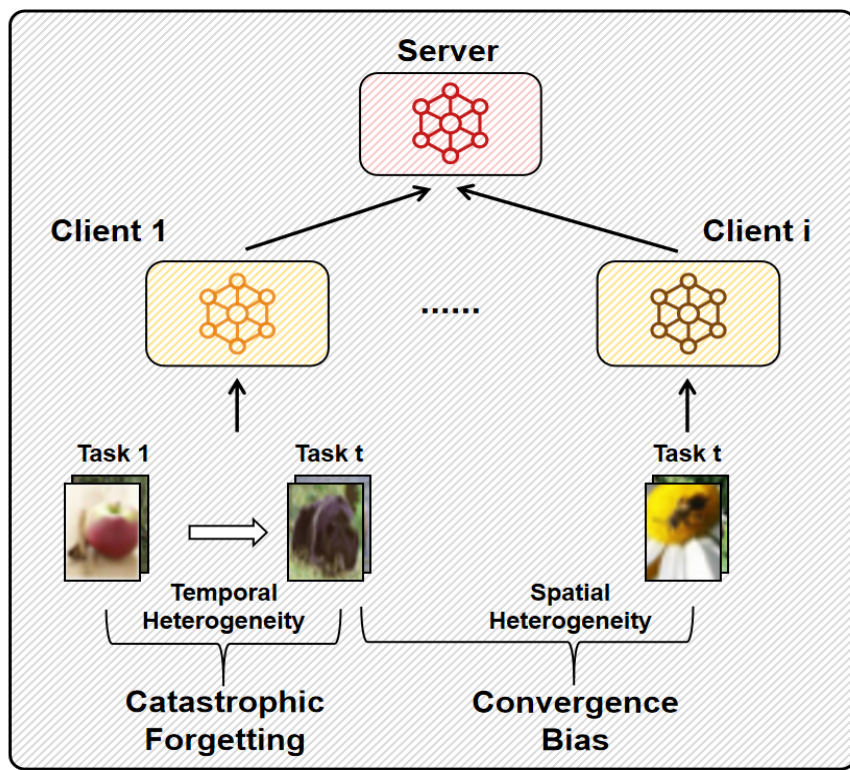
Stable online updates

减弱病态矩阵影响，提升更新稳定性。

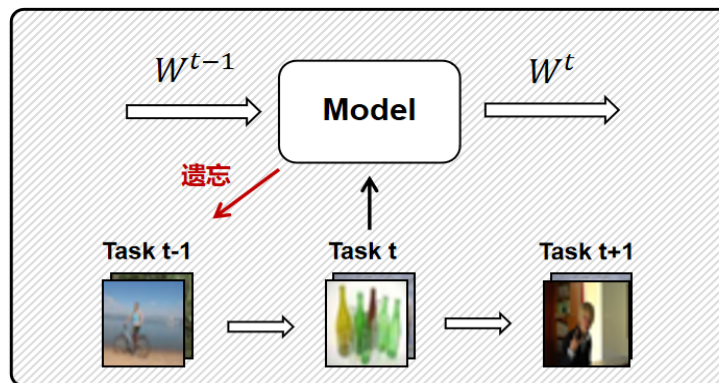
Inverse-free federated learning combines the advantages of single-wheel, efficient, robust, and stable. This further enhances the application capability of broad federated learning in complex scenarios.

A federated learning framework for continuous data flows

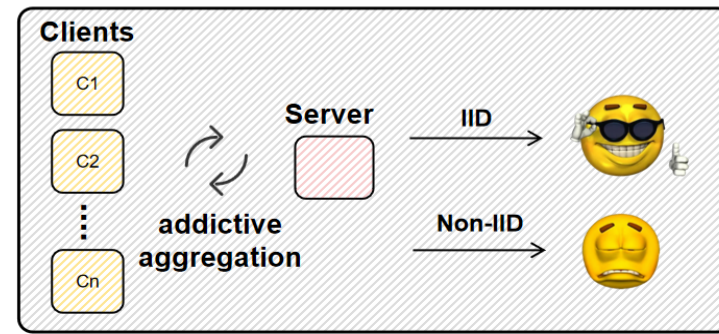
Traditional federated learning usually assumes fixed tasks and struggles to cope with **knowledge forgetting** when new tasks arrive, thus introducing the concept of **continual learning**. Federated continual learning achieves continuous adaptation of models to new knowledge under **privacy protection** and distributed constraints, while striving to retain old knowledge.



(a) 联邦持续学习的两大挑战



(b) 灾难性遗忘



(c) 不一致收敛

难点：数据在时间和空间上的分布都不均衡，数据时间异质性会引发模型对早期知识的**灾难性遗忘**；而数据空间异质性会导致聚合损失曲面发散，引发**不一致收敛**。在多轮学习中，两者还会存在叠加效应，进一步加剧模型性能的退化。



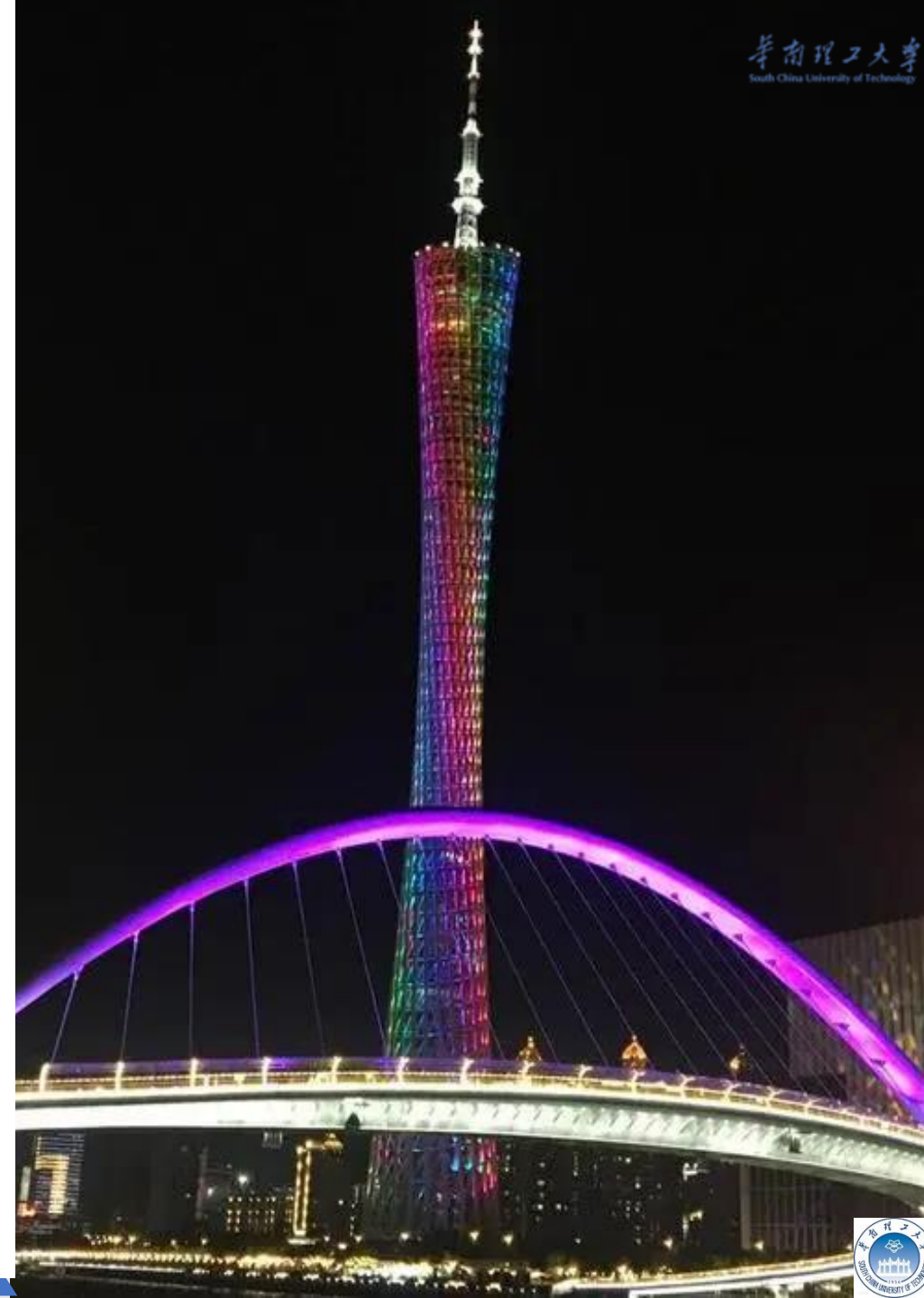
2024 年广州市科普作品大赛获奖作品榜单

2024 年广州市优秀科普作品（图书）

序号	图书名称	主要作者	出版社	推荐单位（推荐人）
1	《“口袋里的人工智能”系列丛书》	陈俊龙（总主编）	广东科技出版社	广东科技出版社
2	《癌症那些事》	万德森	广东科技出版社	广东科技出版社
3	《藏在节气里的中医》	曾科学	中国中医药出版社	广东省第二中医院
4	《常见心血管病调养功法》	彭锐（广州中医药大学第一附属医院）	广东科技出版社	自荐（彭锐）
5	《常用补益中药鉴别与应用》	夏鑫华	广东科技出版社	广州医科大学附属第一医院
6	《读山、听史、观自然 白云山科普丛书》	广州市白云山风景名胜区管理局	花城出版社	广州市白云山风景名胜区管理局、广东花城出版社



Currently, it has exported Mongolian, German, Arabic, Malay, and other foreign language copyrights



**Artificial intelligence
technology empowers the
development of digital
economy industry**

Thank you, everyone!