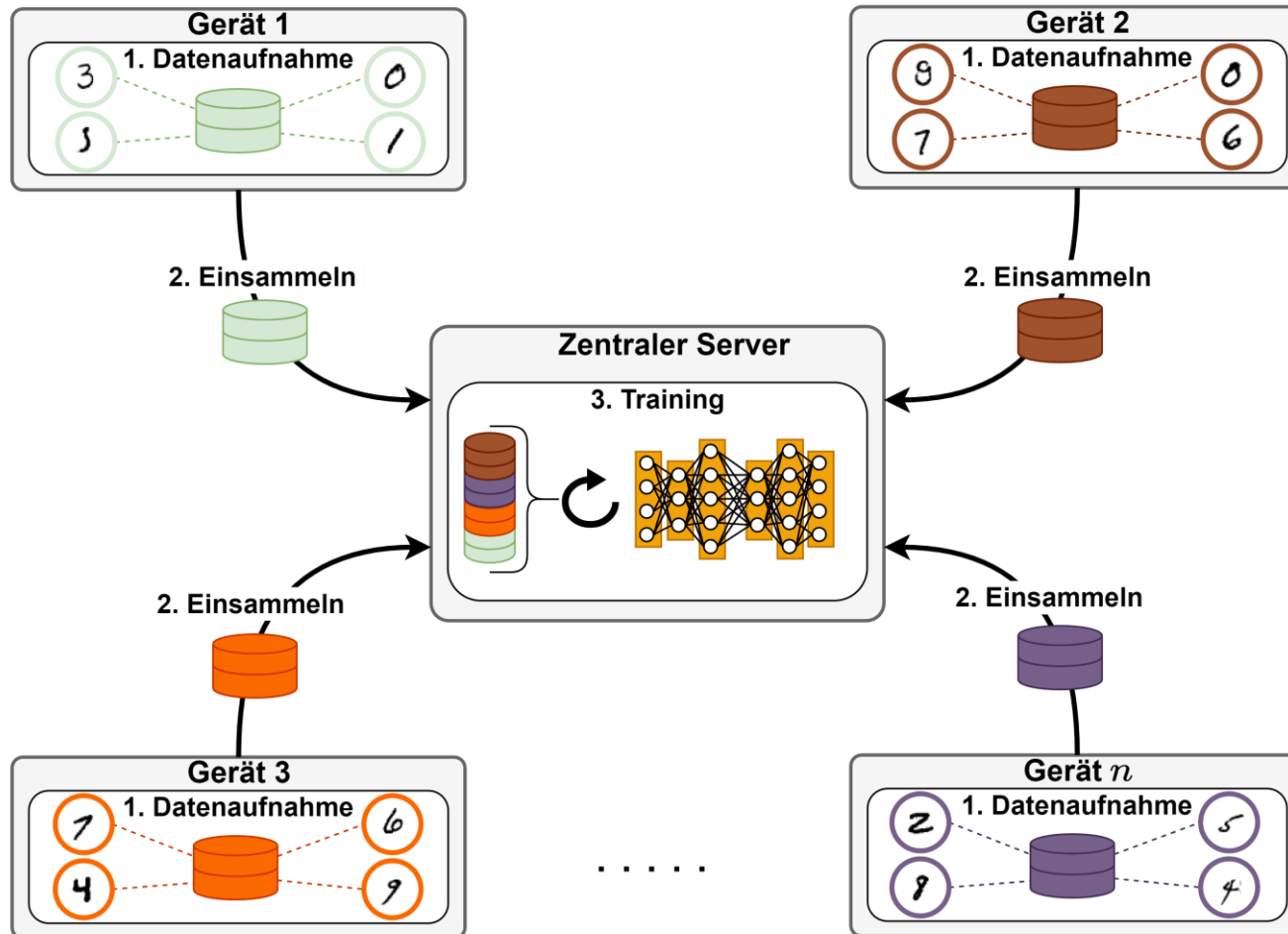


Von Zentral zu Edge: Dezentrales Federated Learning auf Mikrocontrollern

Dr.-Ing. Lars Wulfert
Smart Systems Conference, 19.11.2025, Dortmund

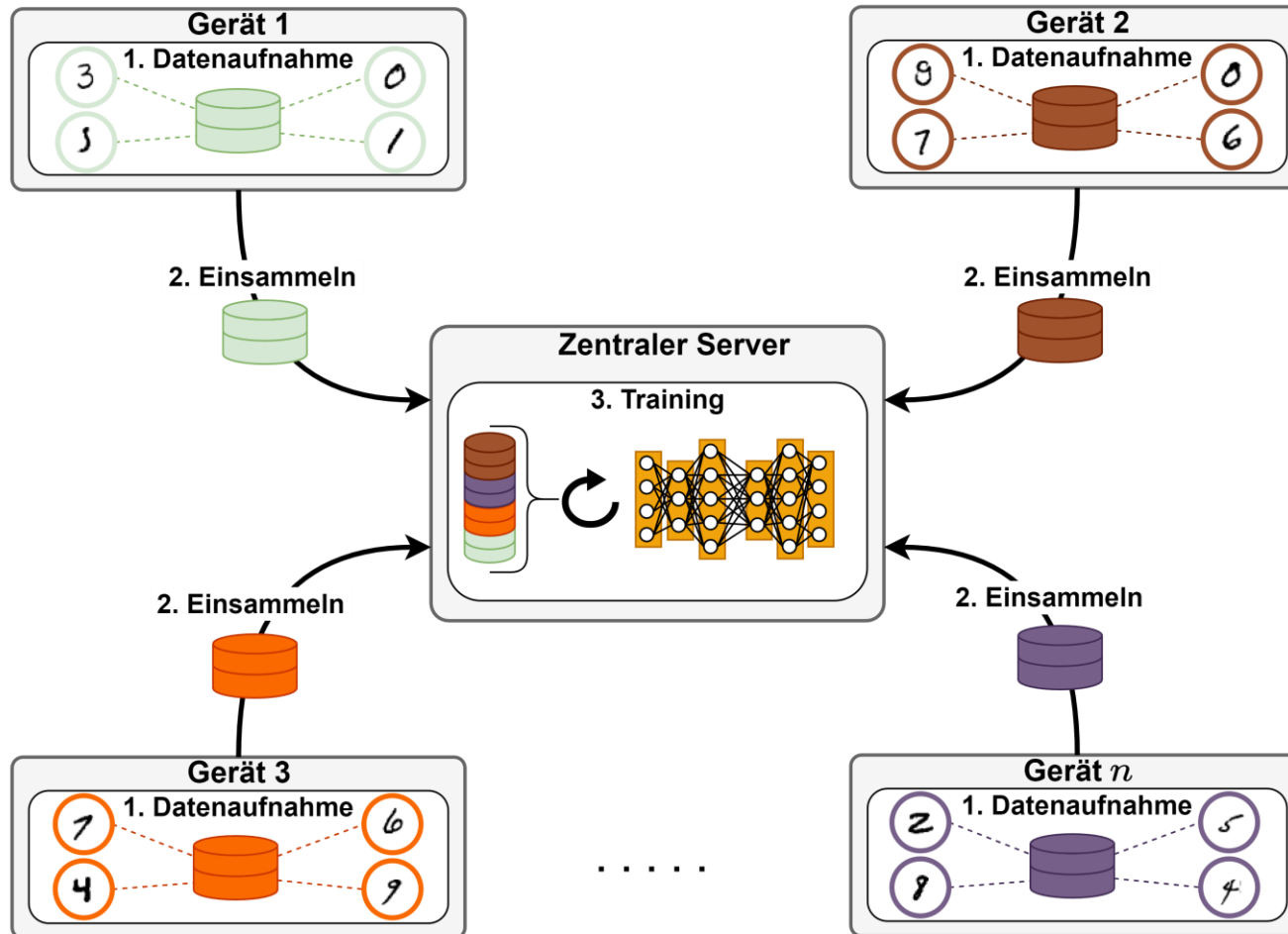
Machine Learning



Maschinelles Lernen (ML)

1. Datenaufnahme auf den Geräten
2. Daten vom Server einsammeln
3. Training des Modells im Server
4. Ausführung des Modells auf dem Gerät

Machine Learning

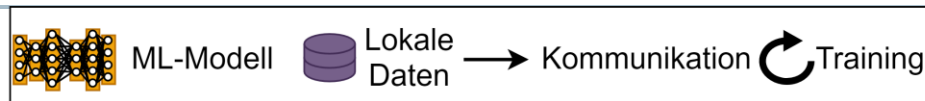


Maschinelles Lernen (ML)

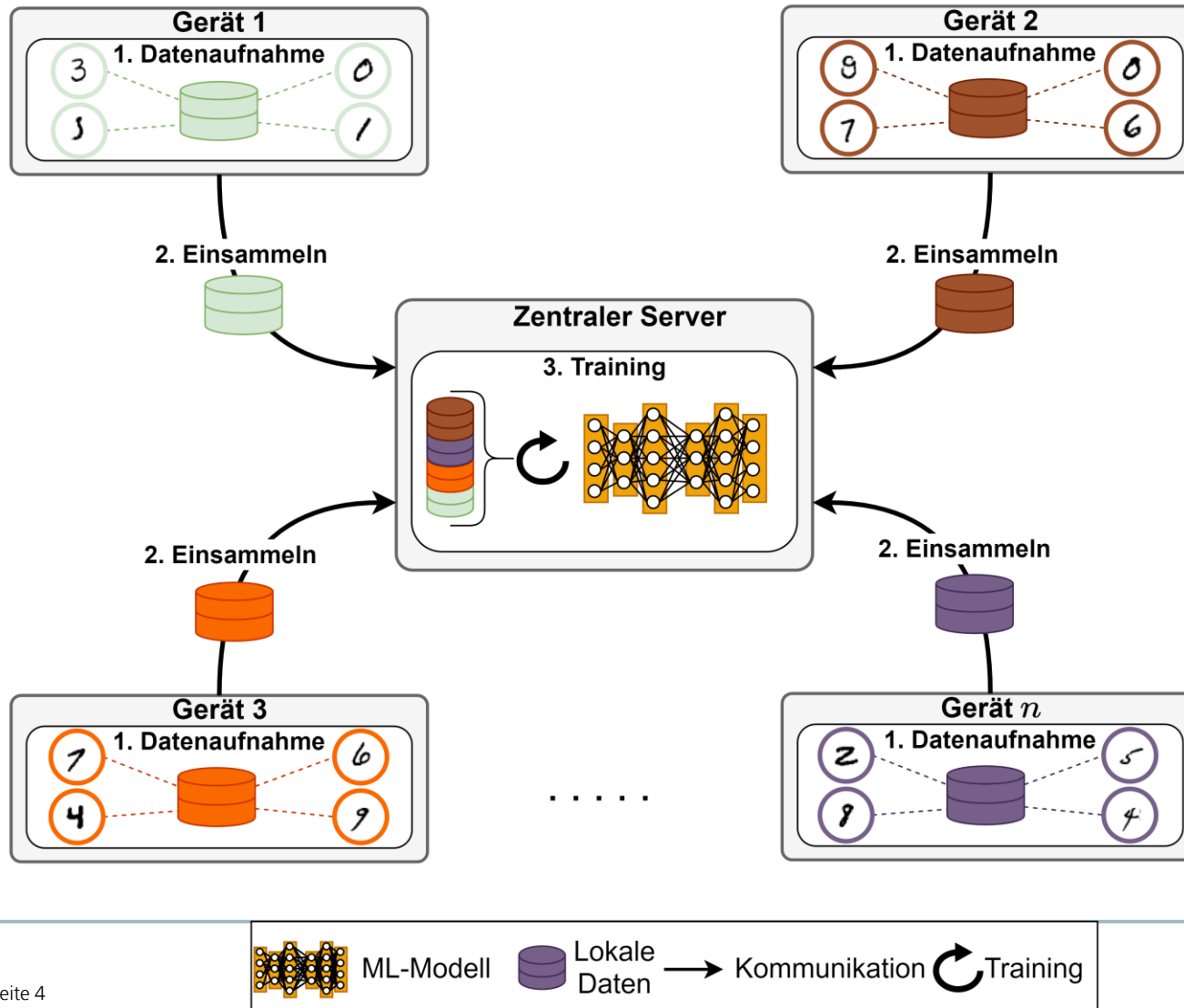
1. Datenaufnahme auf den Geräten
2. Daten vom Server einsammeln
3. Training des Modells im Server
4. Ausführung des Modells auf dem Gerät

Vorteile

- ✓ Verwendung performanter Server
- ✓ Skalierbarkeit
- ✓ Nutzung zentraler Datengrundlage



Machine Learning



Maschinelles Lernen (ML)

1. Datenaufnahme auf den Geräten
2. Daten vom Server einsammeln
3. Training des Modells im Server
4. Ausführung des Modells auf dem Gerät

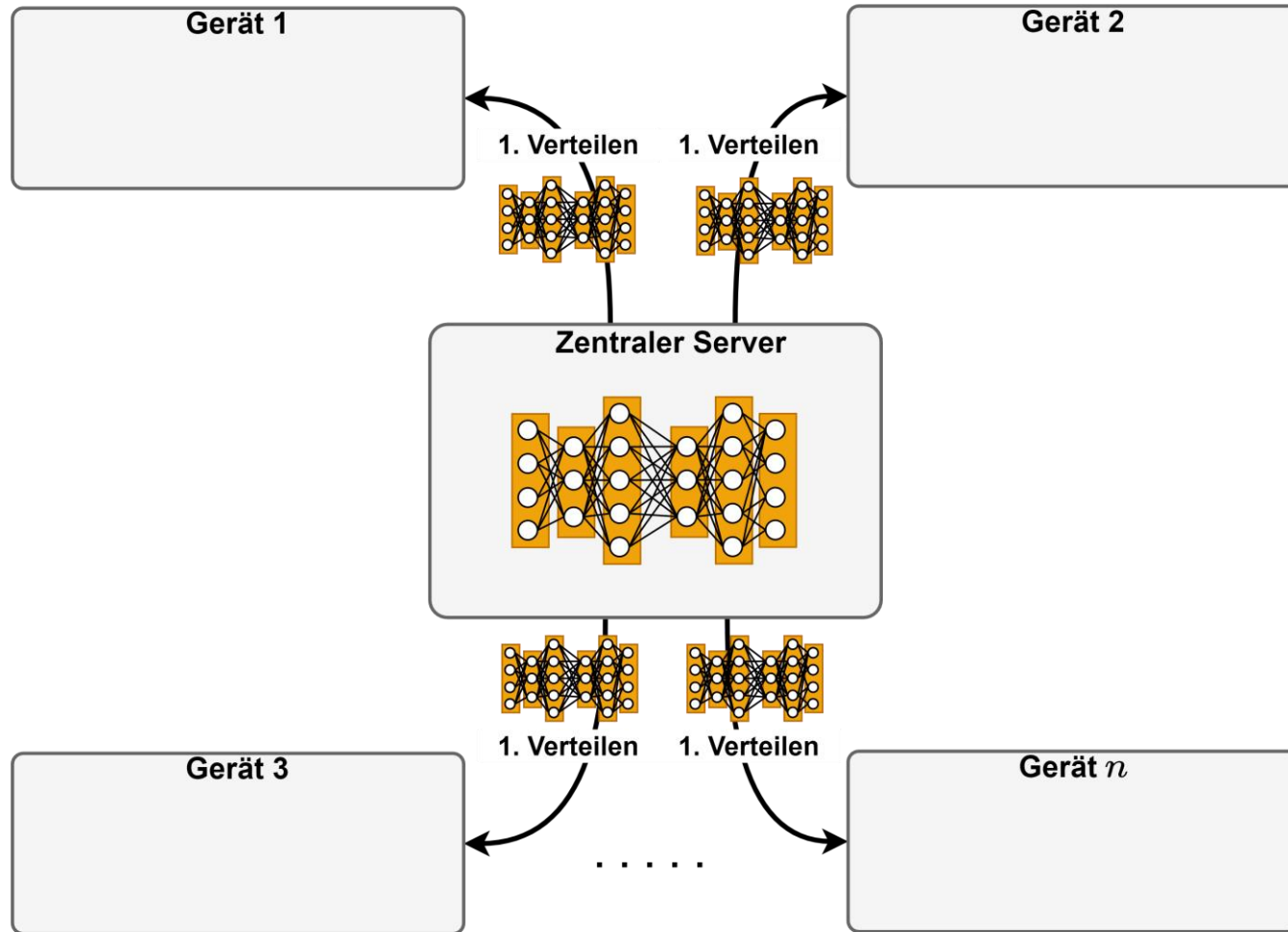
Vorteile

- ✓ Verwendung performanter Server
- ✓ Skalierbarkeit
- ✓ Nutzung zentraler Datengrundlage

Nachteile

- × Netzwerkbelastung
- × DSGVO kann die Verwendung von personenbezogenen Daten einschränken
- × **Datensicherheit & Privatsphäre von sensiblen Daten nicht immer gewährleistet**

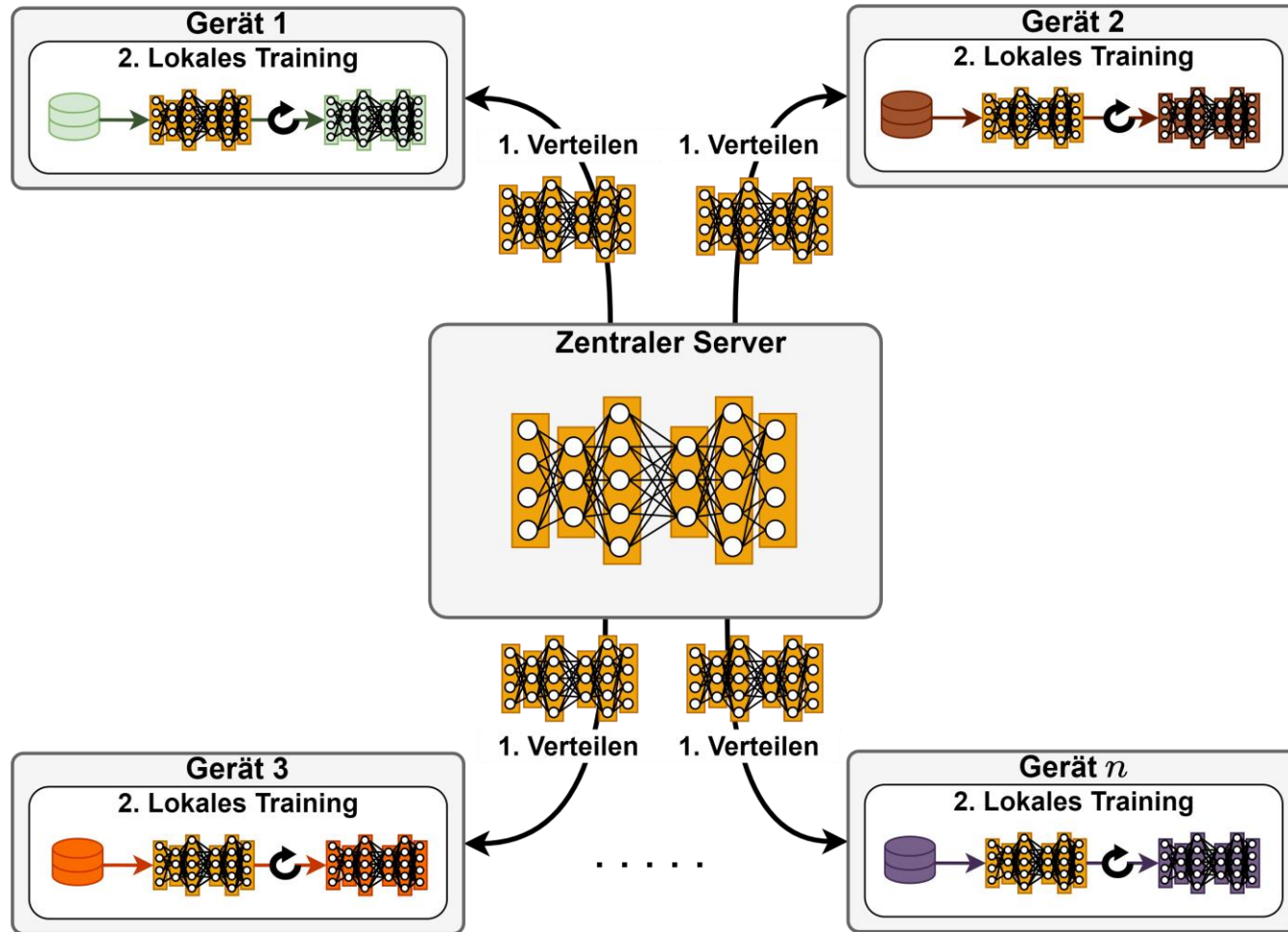
Zentrales Federated Learning



Zentrales föderales Lernen (ZFL)

1. Globales Modell verteilen

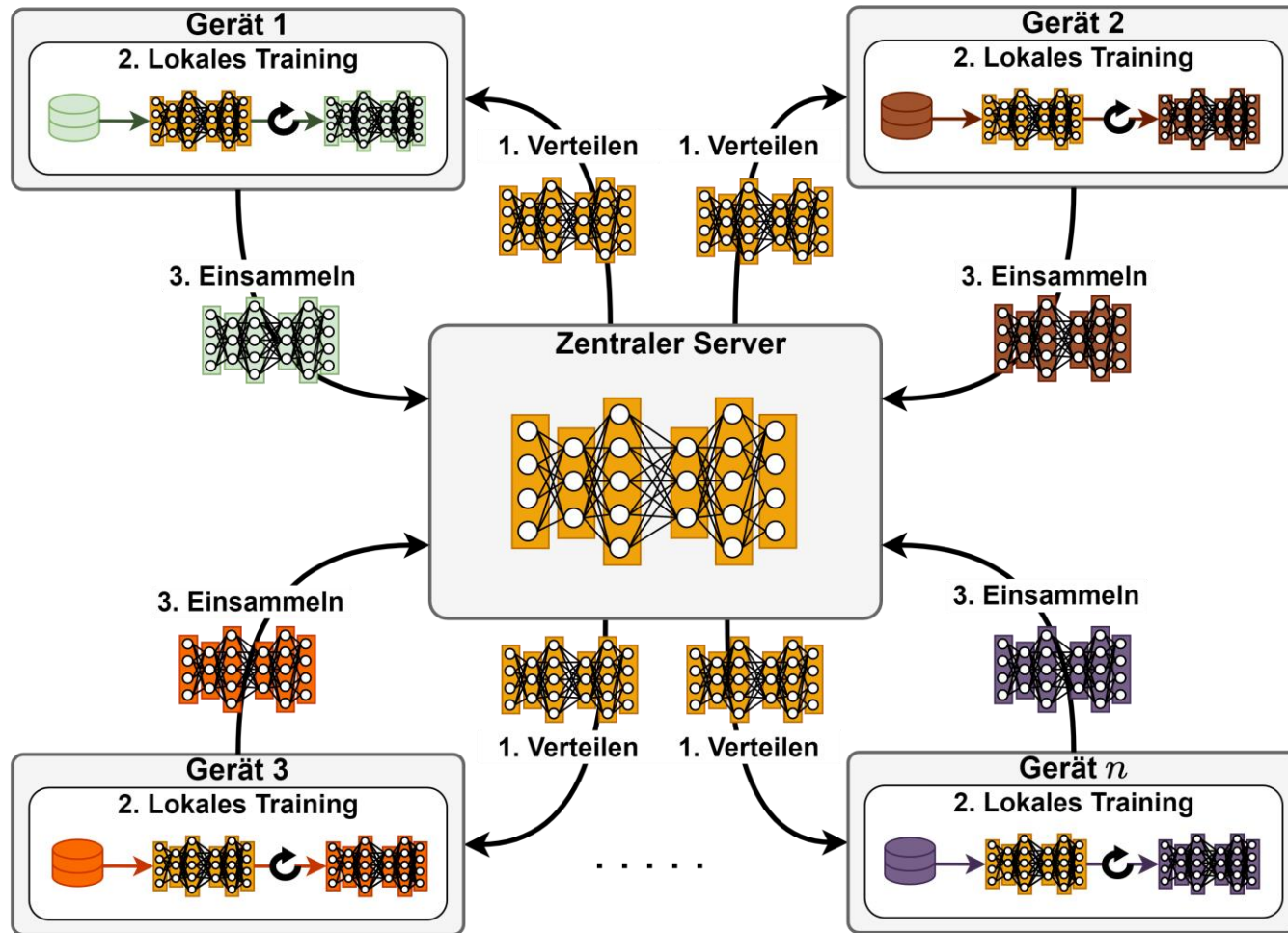
Zentrales Federated Learning



Zentrales föderales Lernen (ZFL) [12]

1. Globales Modell verteilen
2. Lokales Training

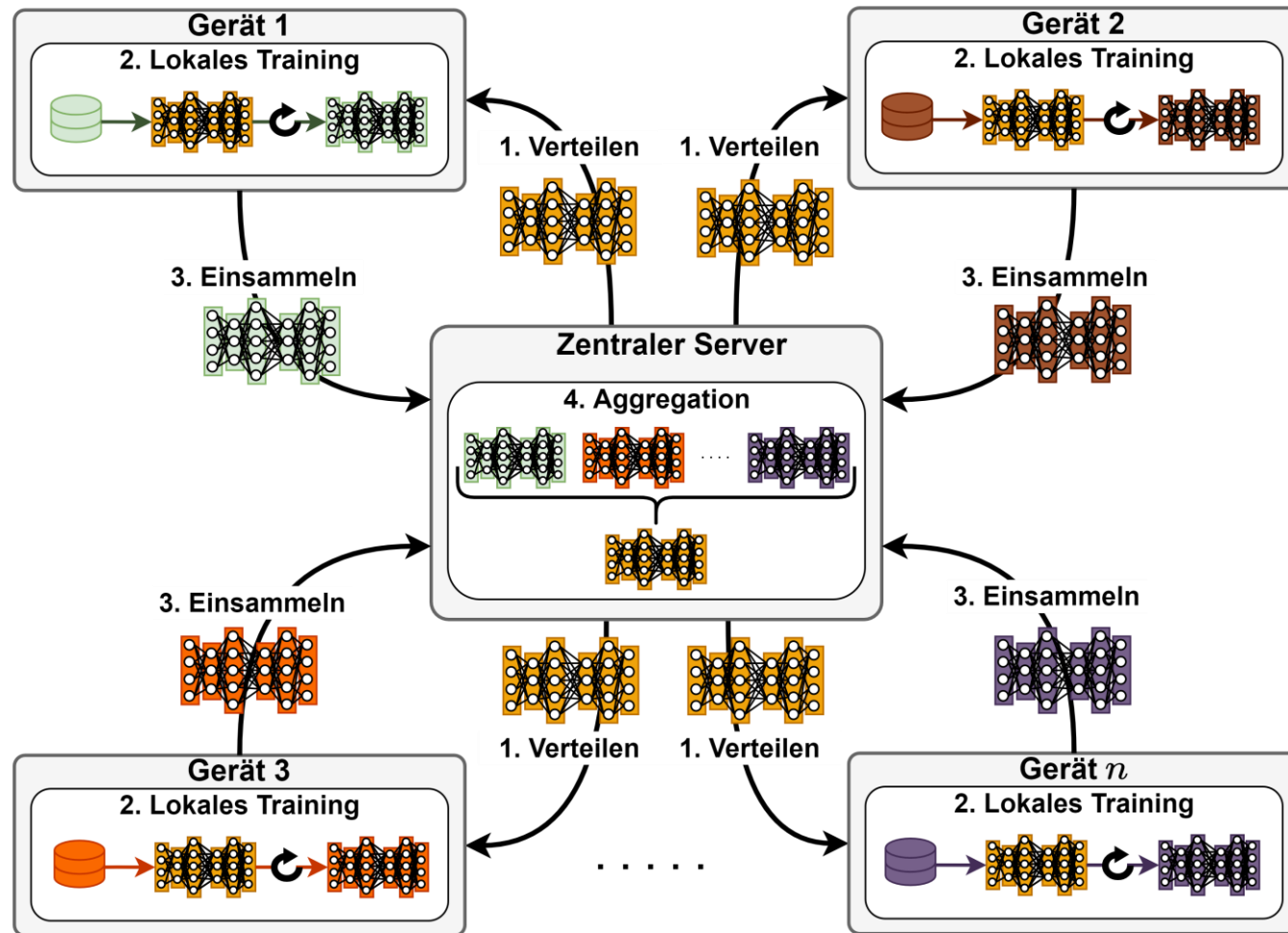
Zentrales Federated Learning



Zentrales föderales Lernen (ZFL) [12]

1. Globales Modell verteilen
2. Lokales Training
3. Einsammeln der lokalen Modelle

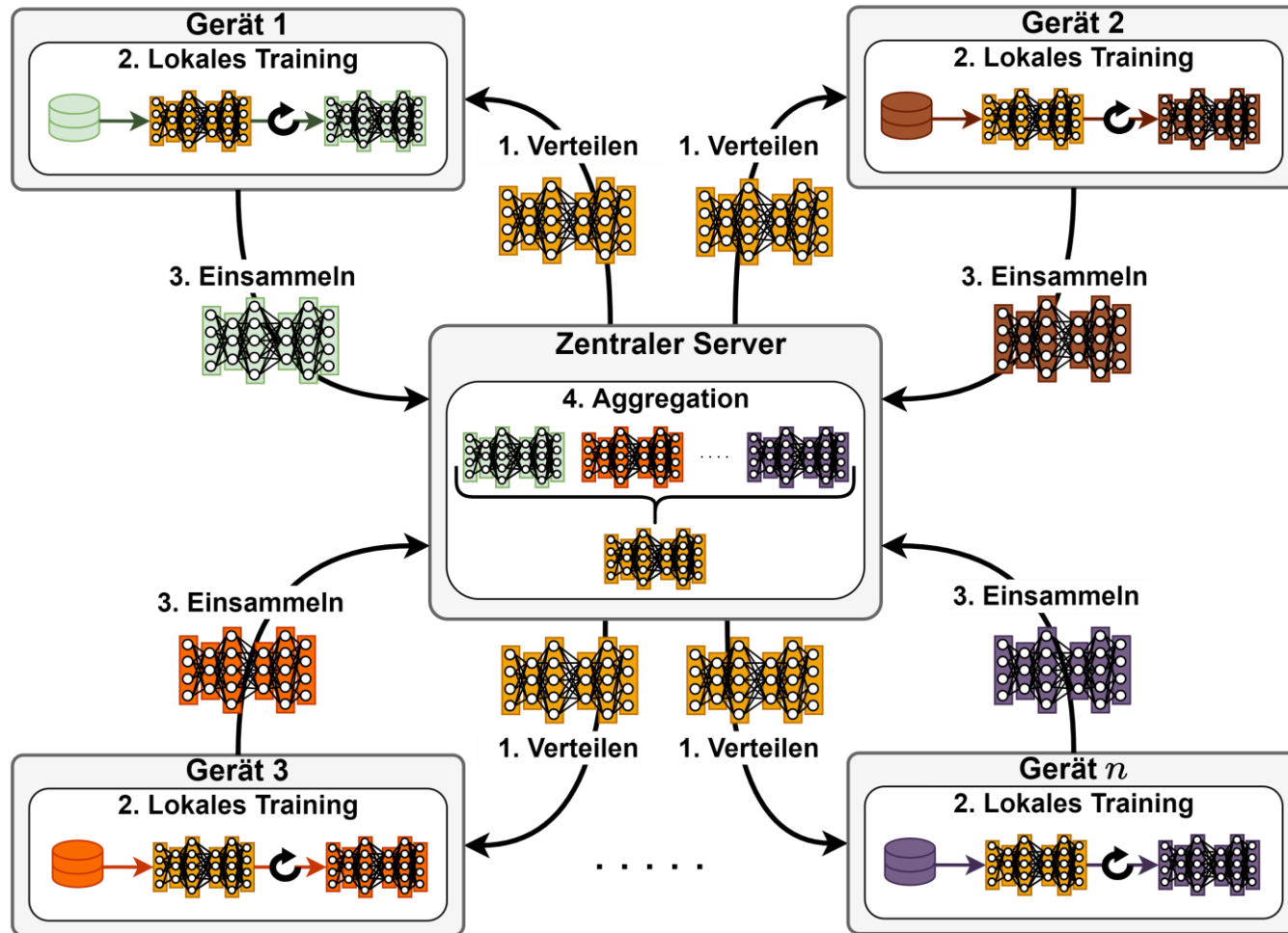
Zentrales Federated Learning



Zentrales föderales Lernen (ZFL)

1. Globales Modell verteilen
2. Lokales Training
3. Einsammeln der lokalen Modelle
4. Aggregation

Zentrales Federated Learning



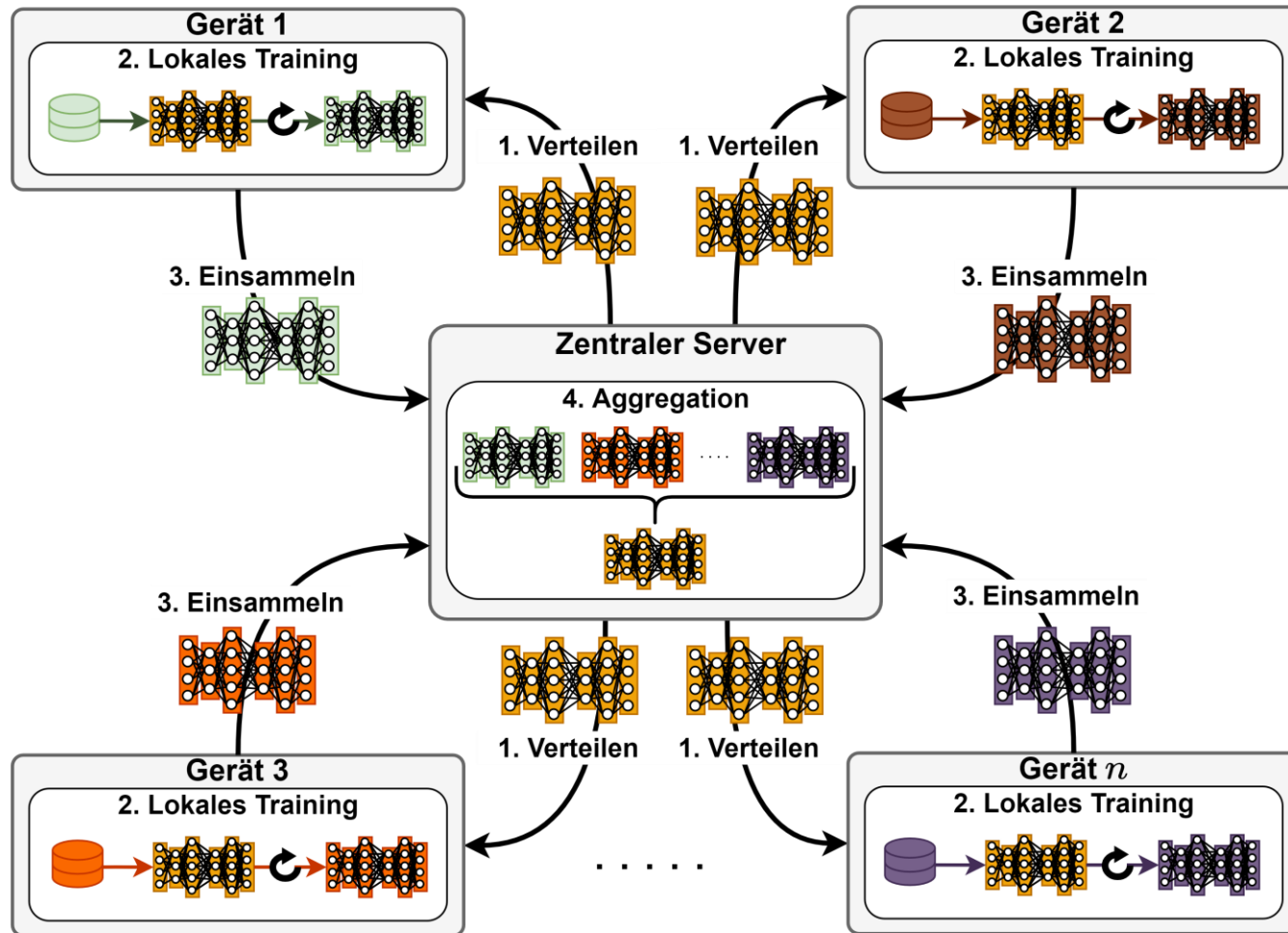
Zentrales föderales Lernen (ZFL)

1. Globales Modell verteilen
2. Lokales Training
3. Einsammeln der lokalen Modelle
4. Aggregation

Vorteile

- ✓ Neue Daten werden schnell integriert
- Reduzierung der Datenübertragung
- ✓ Verbesserung der Datensicherheit

Zentrales Federated Learning



Zentrales föderales Lernen (ZFL)

1. Globales Modell verteilen
2. Lokales Training
3. Einsammeln der lokalen Modelle
4. Aggregation

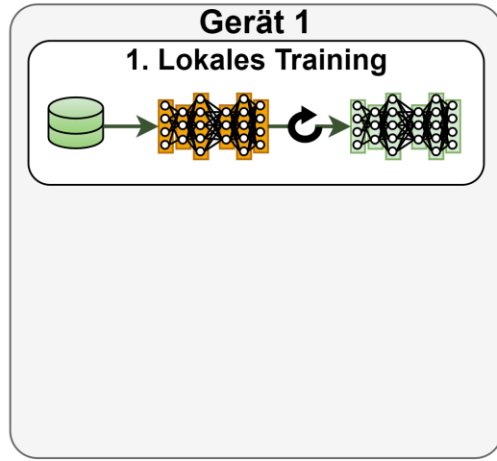
Vorteile

- ✓ Neue Daten werden schnell integriert
- ✓ Reduzierung der Datenübertragung
- ✓ Verbesserung der Datensicherheit

Nachteile

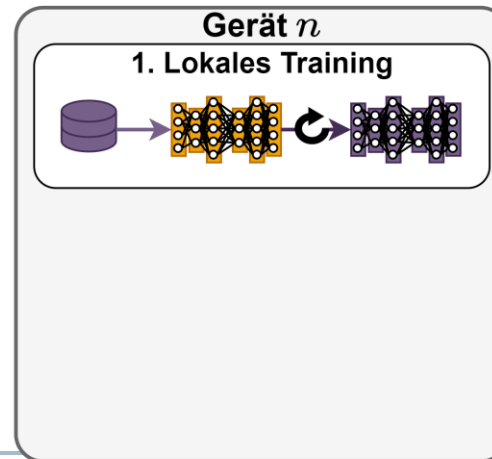
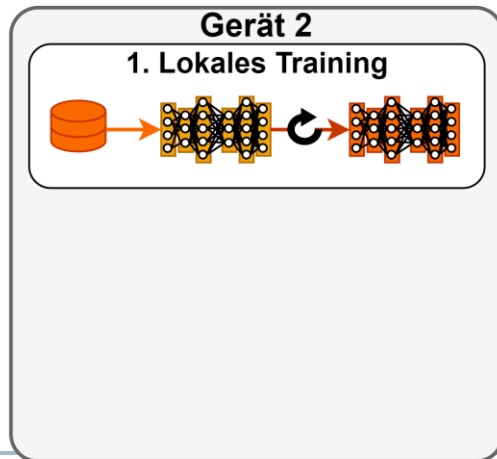
- ✗ Ausfall des zentralen Servers kann ZFL verhindern
- ✗ Schlechte Skalierbarkeit
- ✗ Server kann korrumpiert sein

Dezentrales Federated Learning

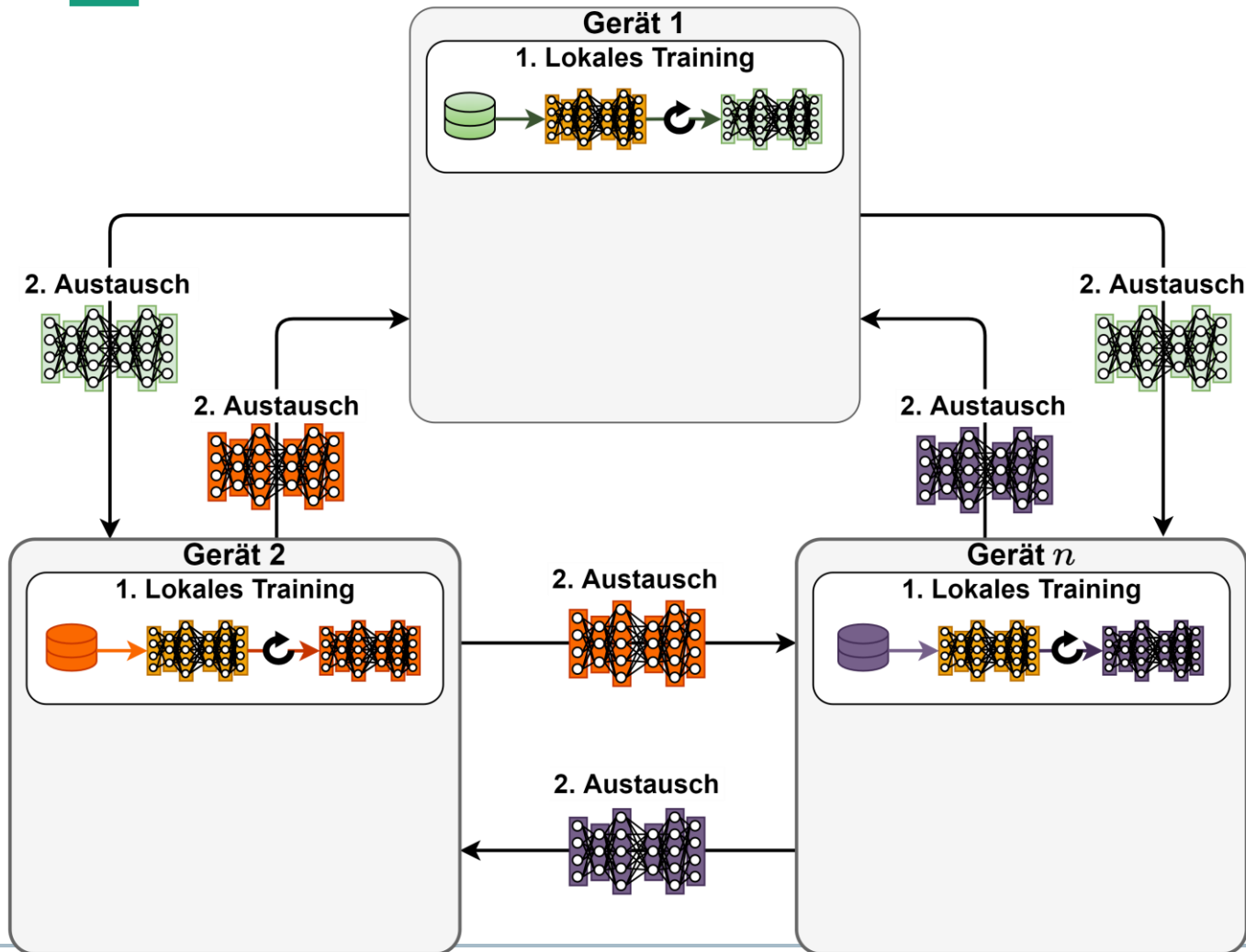


Dezentrales föderales Lernen (DFL)

1. Lokales Modell Training



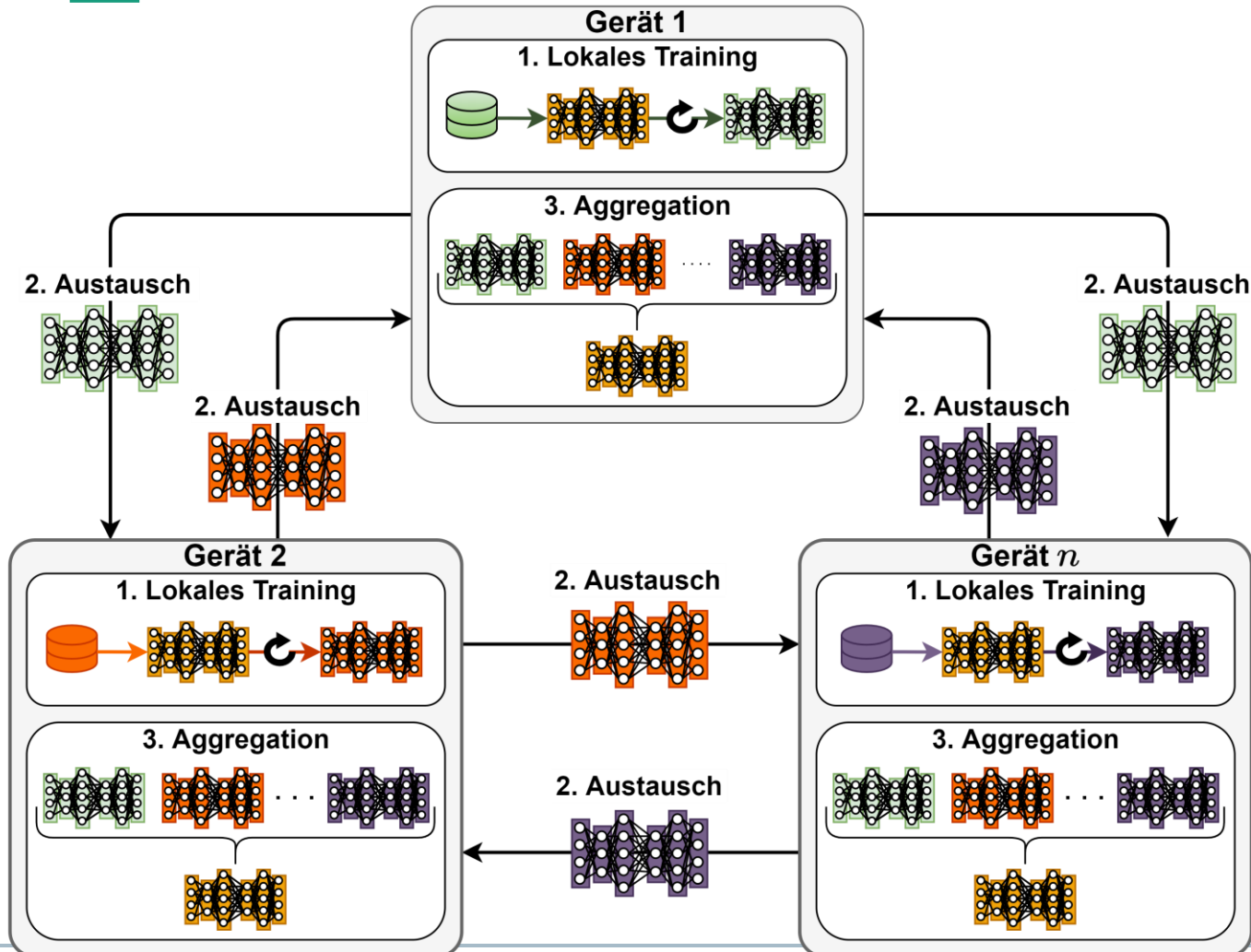
Dezentrales Federated Learning



Dezentrales föderales Lernen (DFL)

1. Lokales Modell Training
2. Modelle austauschen

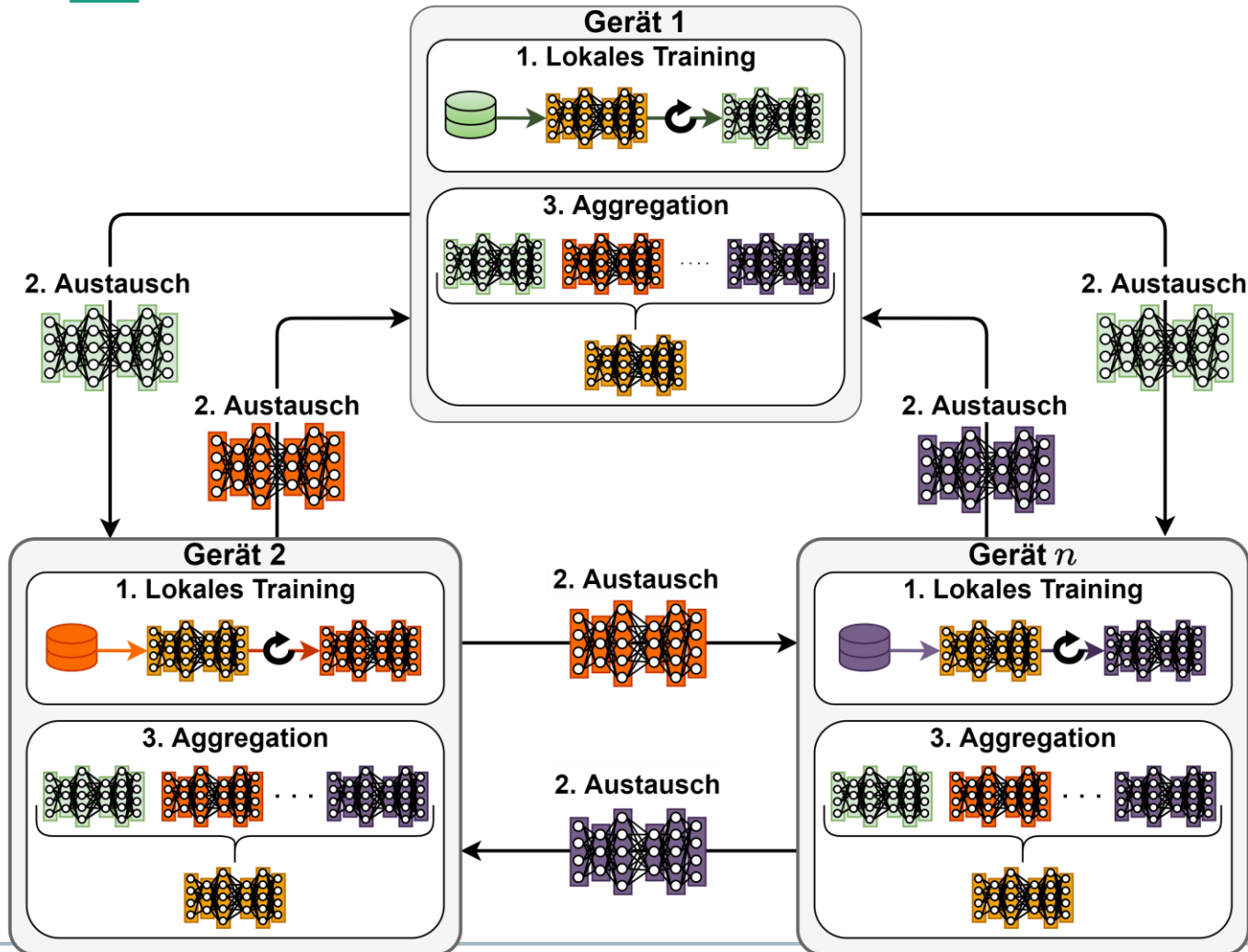
Dezentrales Federated Learning



Dezentrales föderales Lernen (DFL)

1. Lokales Modell Training
2. Modelle austauschen
3. Aggregation

Dezentrales Federated Learning



Dezentrales föderales Lernen (DFL)

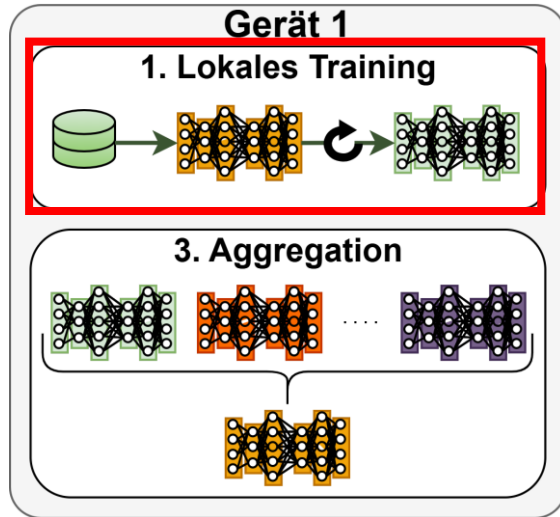
1. Lokales Modell Training
2. Modelle austauschen
3. Aggregation

Nachteile & Herausforderungen

- × Eine heterogene Datenverteilungen stellt Herausforderungen an die Genauigkeit
- × Hoher Kommunikationsaufwand
- × On-Device Training auf Mikrocontrollern stellt Herausforderung an die Speicherkapazität

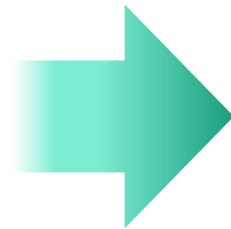
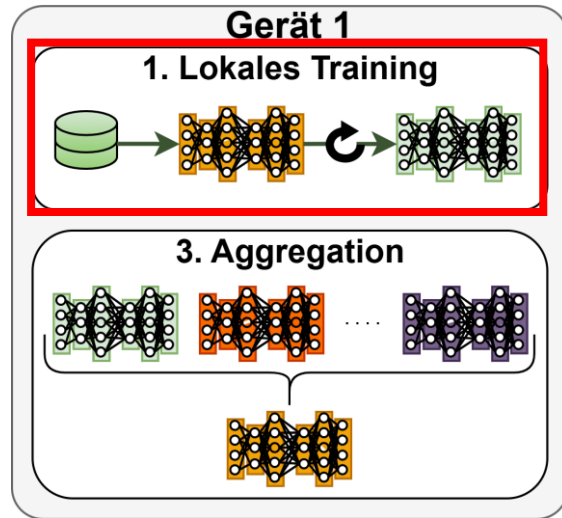
Herausforderung 1

On-Device Training



Herausforderung 1

On-Device Training



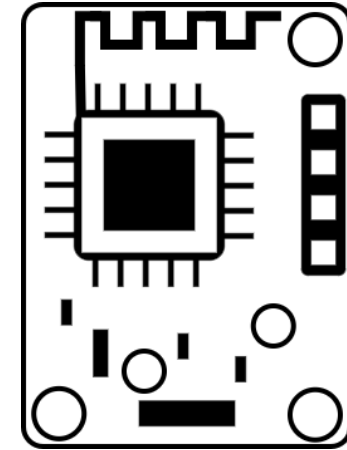
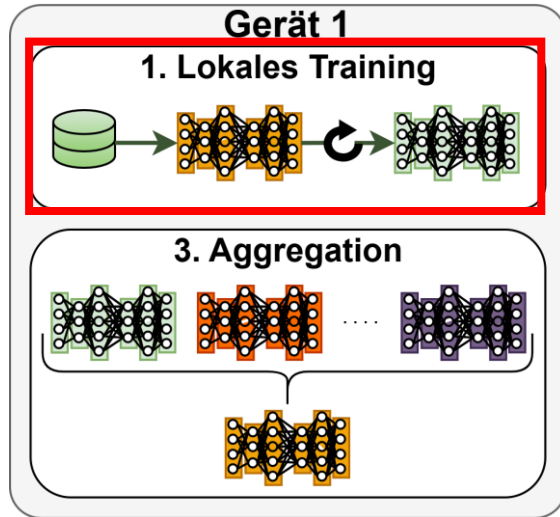
 PyTorch

 TensorFlow

 Keras

Herausforderung 1

On-Device Training

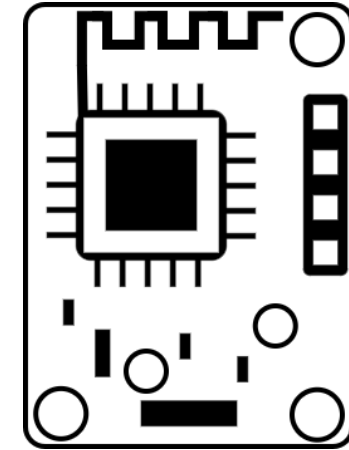
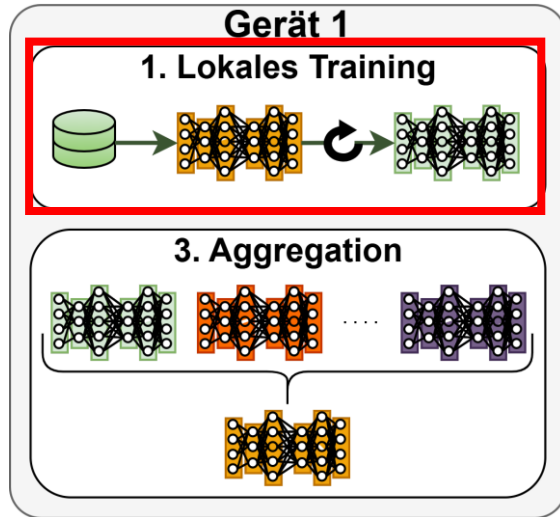


Speicher	Cloud ML (NVIDIA V100 [34])	<u>Red.</u> →	Mobile ML (iPhone 15 [35])	<u>Red.</u> →	TinyML (ESP32 [36])
RAM	> 16 GB	<u>2,67 ×</u>	> 6 GB	<u>11.538 ×</u>	520 kB

Herausforderung 1

On-Device Training

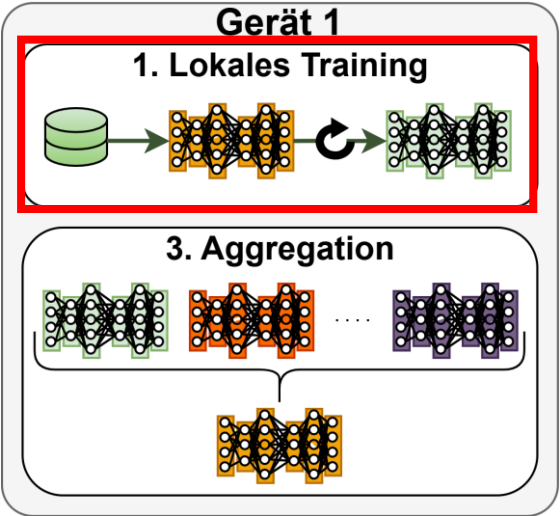
- Frameworks benötigen Python
- Frameworks benötigen zu viel Speicherplatz
- MCU besitzen keine performanten CPUs & GPUs



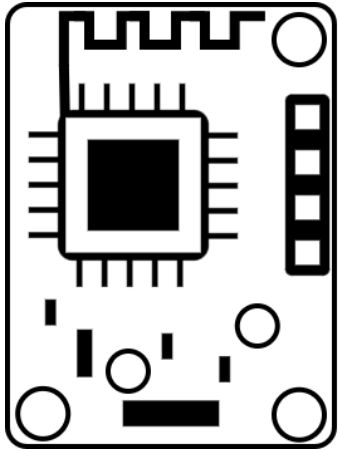
Speicher	Cloud ML (NVIDIA V100 [34])	<u>Red.</u> →	Mobile ML (iPhone 15 [35])	<u>Red.</u> →	TinyML (ESP32 [36])
RAM	> 16 GB	<u>2,67 ×</u>	> 6 GB	<u>11.538 ×</u>	520 kB

Herausforderung 1

On-Device Training



- Frameworks benötigen Python
- Frameworks benötigen zu viel Speicherplatz
- MCU besitzen keine performanten CPUs & GPUs



Speicher	Cloud ML (NVIDIA V100 [34])	Red. →	Mobile ML (iPhone 15 [35])	Red. →	TinyML (ESP32 [36])
RAM	> 16 GB	2,67 × →	> 6 GB	11.538 × →	520 kB

Herausforderung 1

On-Device Training

- Frameworks benötigen Python
- Frameworks benötigen zu viel Speicherplatz
- MCU besitzen keine performanten CPUs & GPUs



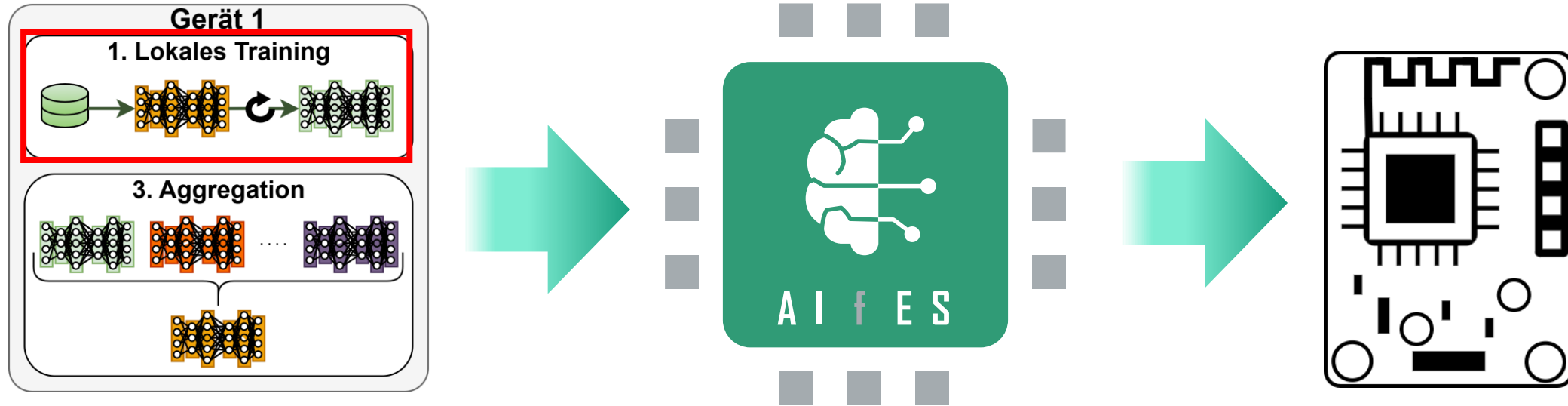
Wie können Machine Learning Modelle auf einem Mikrocontroller trainiert werden?



Speicher	Cloud ML (NVIDIA V100 [34])	Red. →	Mobile ML (iPhone 15 [35])	Red. →	TinyML (ESP32 [36])
RAM	> 16 GB	2,67 ×	> 6 GB	11.538 ×	520 kB

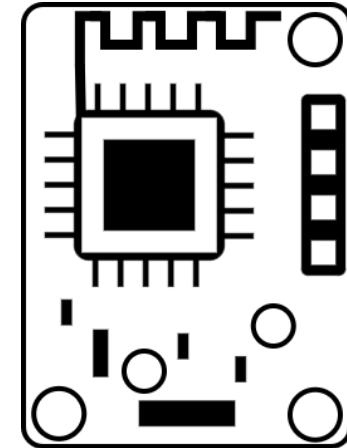
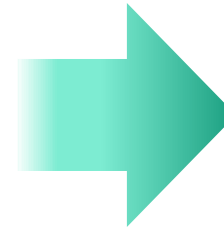
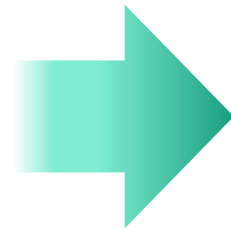
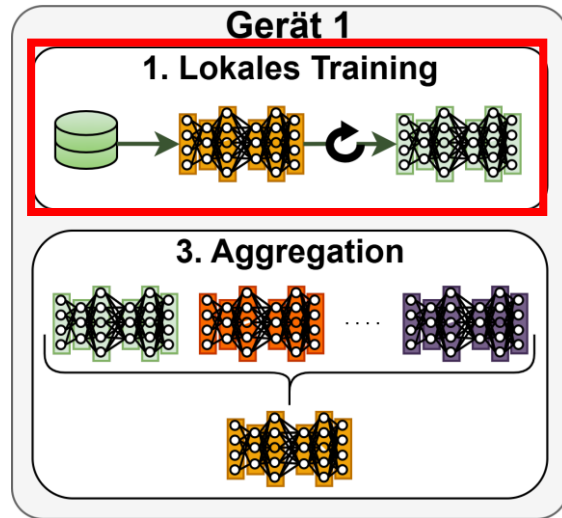
Herausforderung 1

On-Device Training



Herausforderung 1

On-Device Training



Wulfert et al.,
IEEE - TPAMI

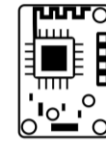
GitHub



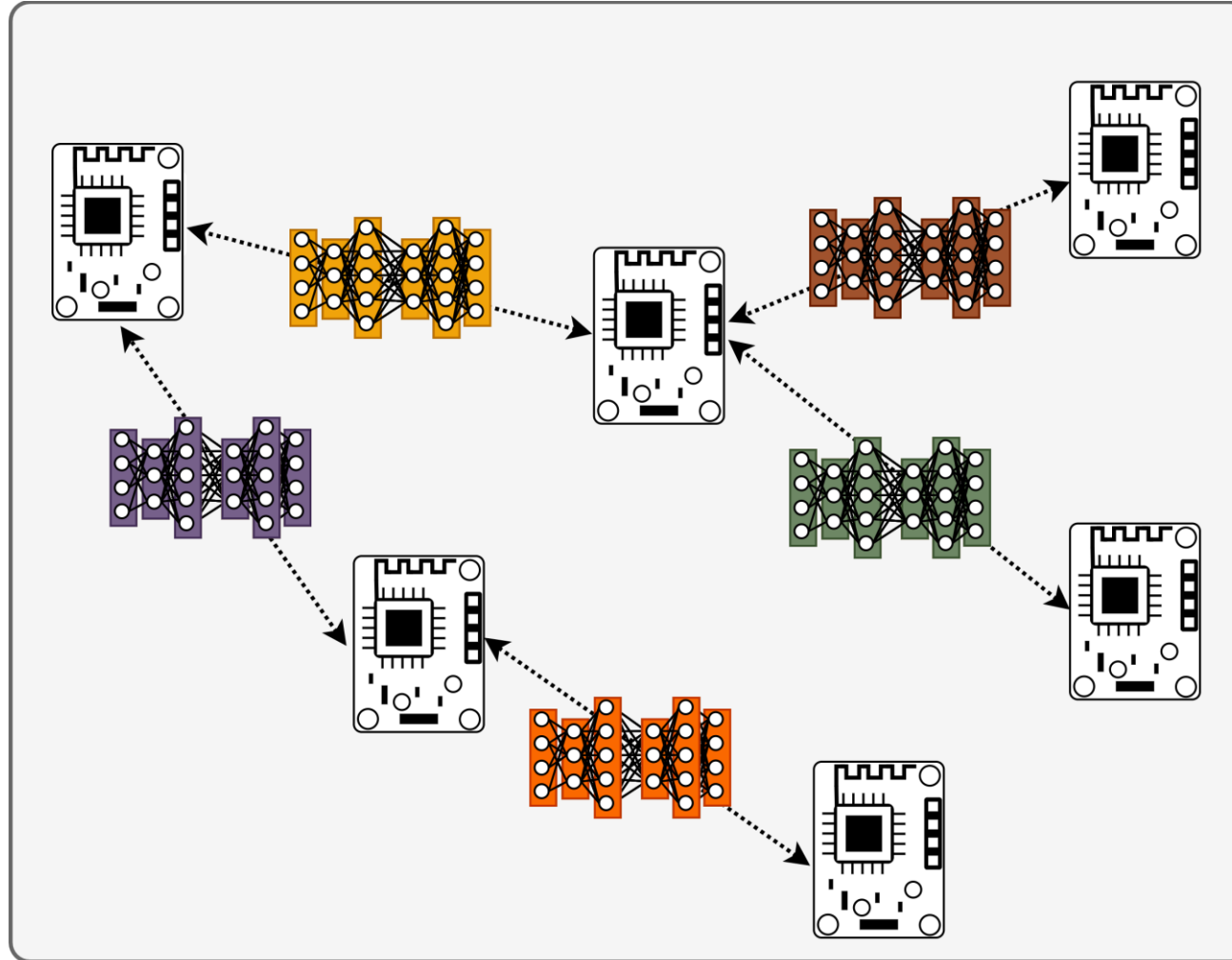
- **AIFES** (AI for Embedded Systems)
- ML Framework für Mikrocontroller und Open Source
- Unterstützt Inference und On-Device Training von
 - Fully Connected Neural Networks
 - Convolutional Neural Networks

Herausforderung 2

Federated Learning auf Mikrocontrollern

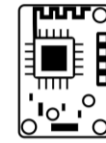


Mikrocontroller $\longleftrightarrow$ Kommunikation

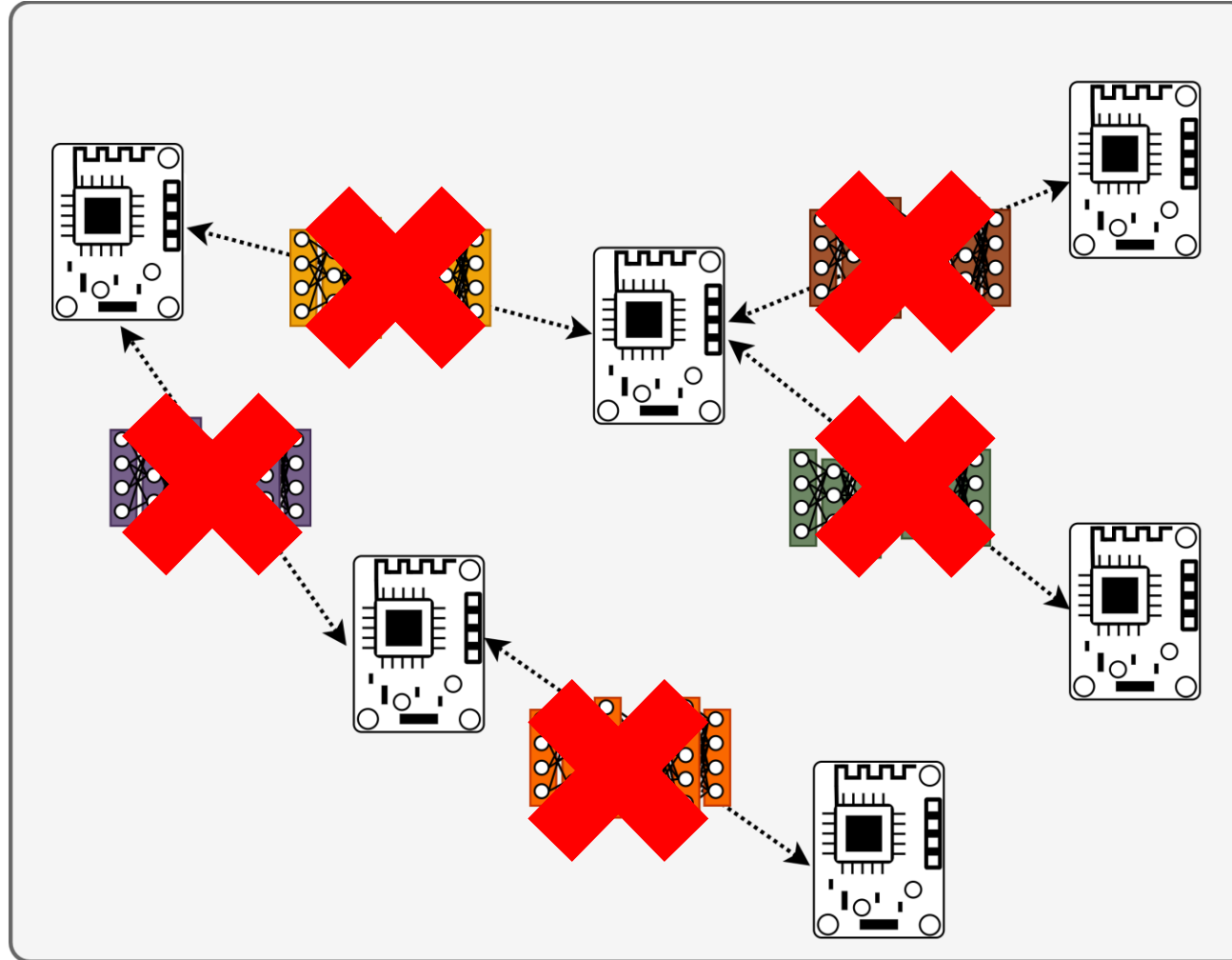


Herausforderung 2

Federated Learning auf Mikrocontrollern

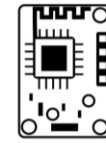


Mikrocontroller $\longleftrightarrow$ Kommunikation



Herausforderung 2

Federated Learning auf Mikrocontrollern



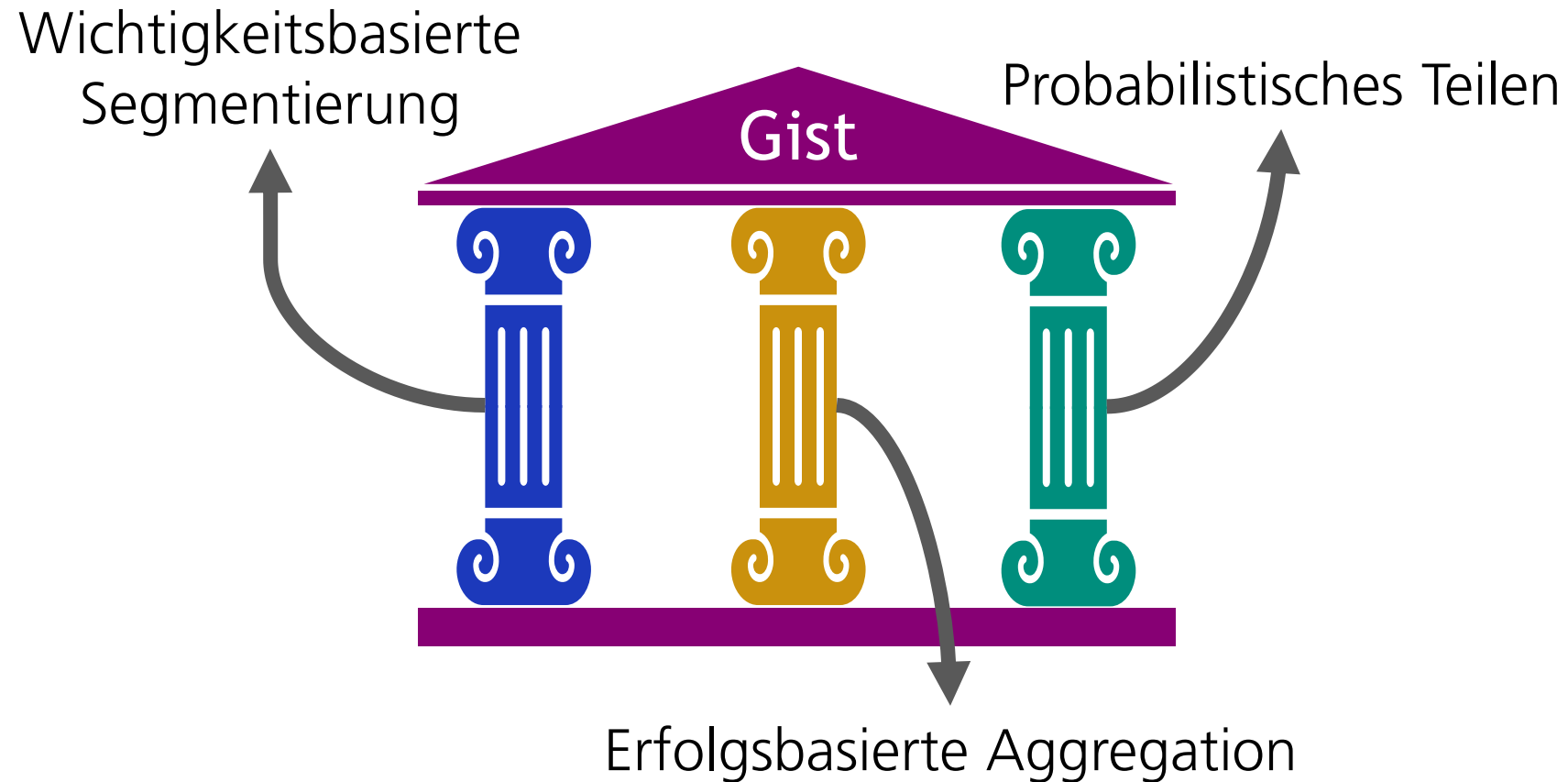
Mikrocontroller $\longleftrightarrow$ Kommunikation

**Ist es möglich, die
Kommunikation/Rechenleistung zu reduzieren
und gleichzeitig die Konvergenz zu erhalten?**



Herausforderung 2

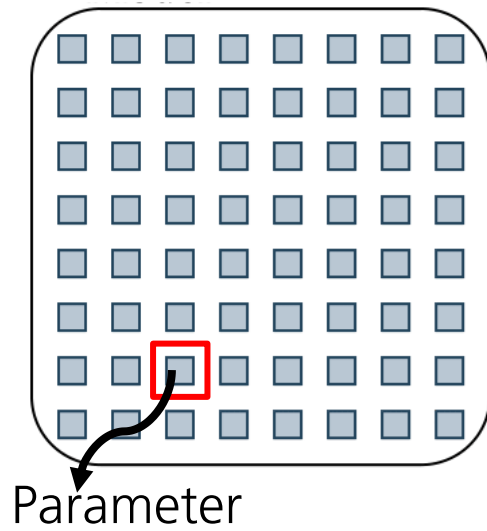
Gist



Gist

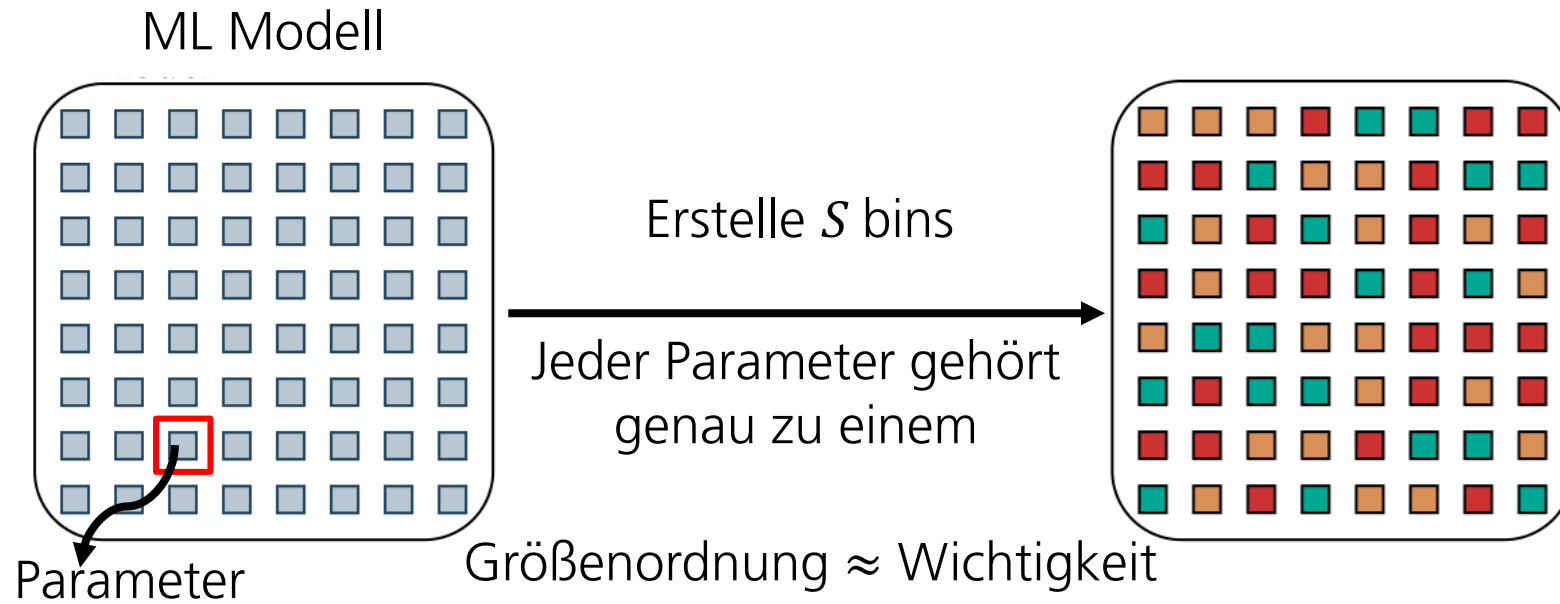
Segmentierung

ML Modell



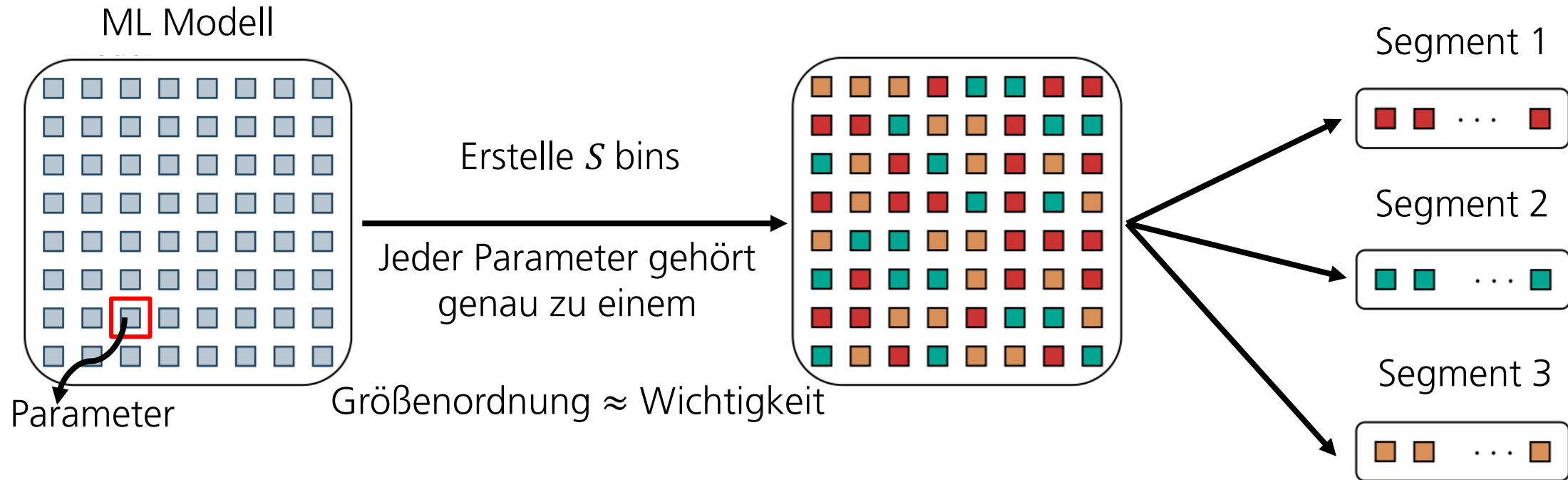
Gist

Segmentierung



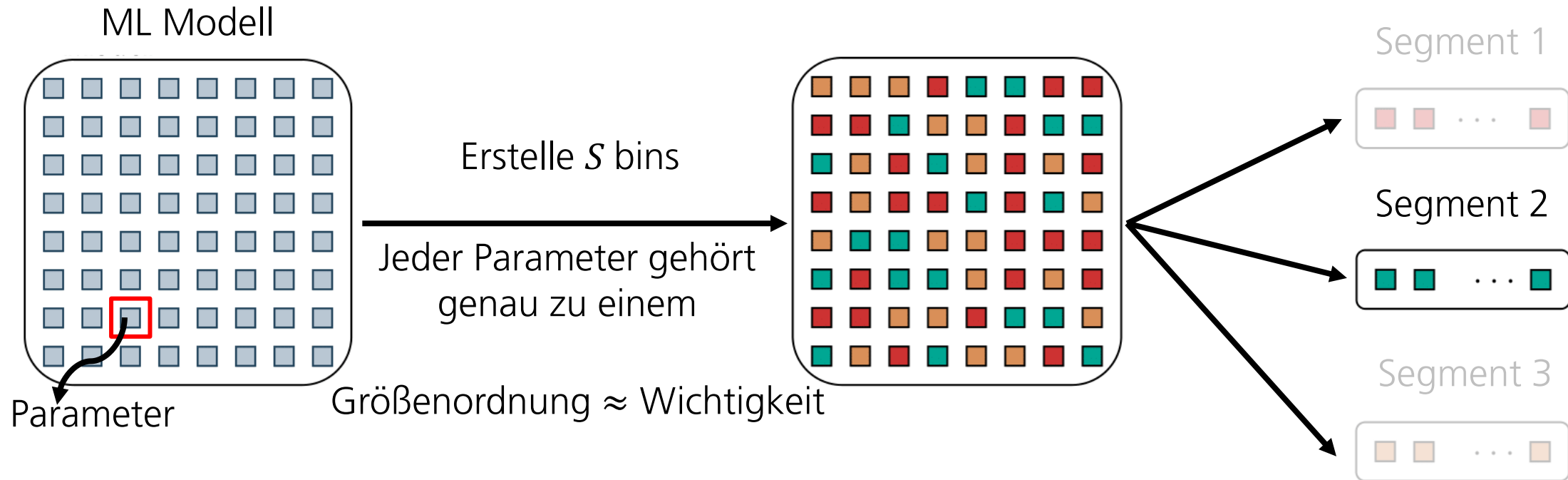
Gist

Segmentierung



Gist

Segmentierung

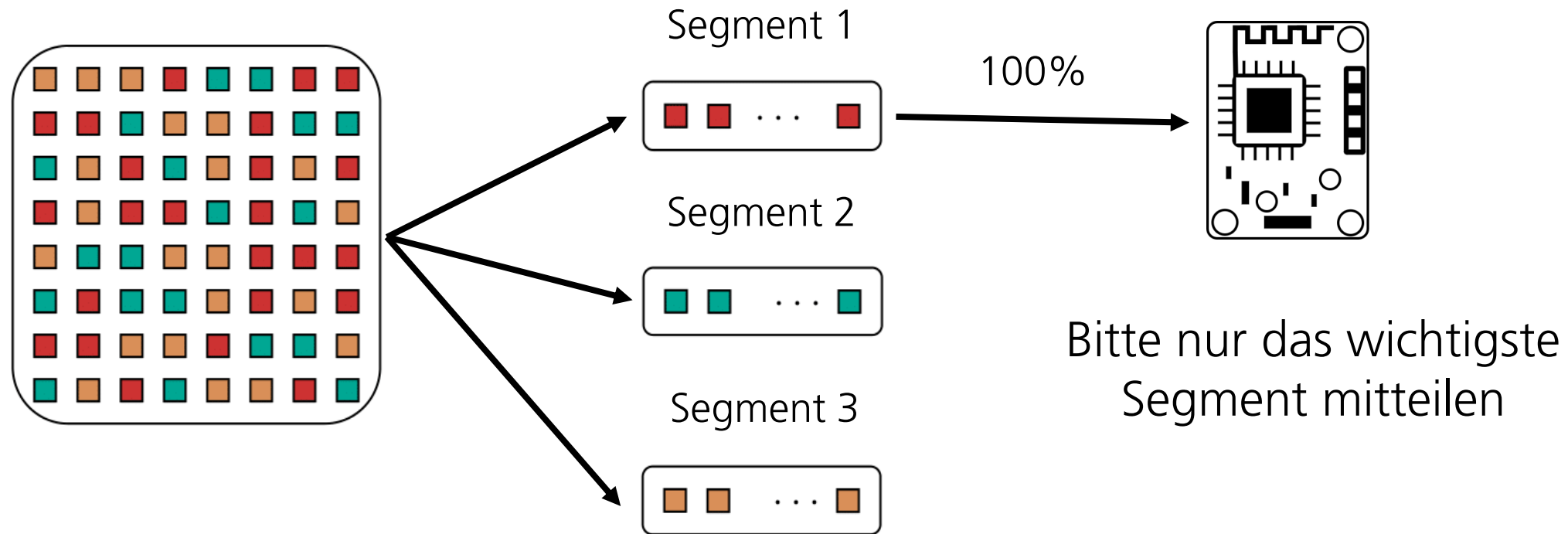


Konzentration des Kommunikationsbudgets auf wirkungsvolle Parameter!



Gist

Probabilistisches Teilen



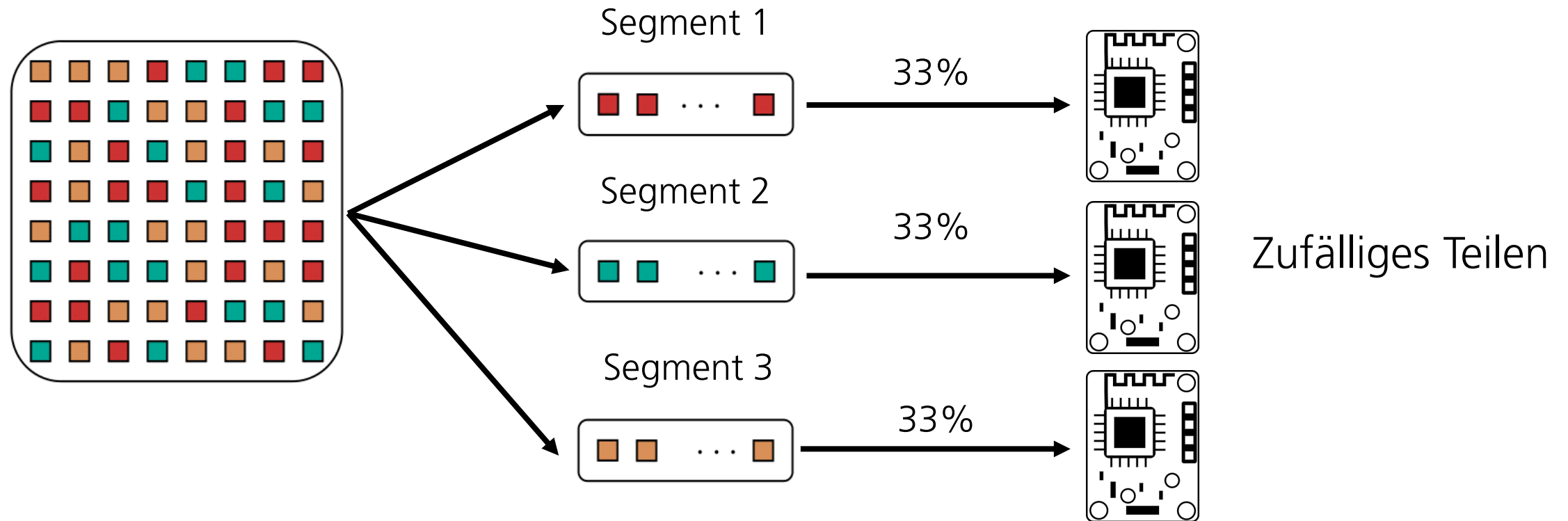
Unzureichende Parameterabdeckung/-vielfalt

→ **Modellabweichung**

→ **Beeinträchtigung der Konvergenz!**

Gist

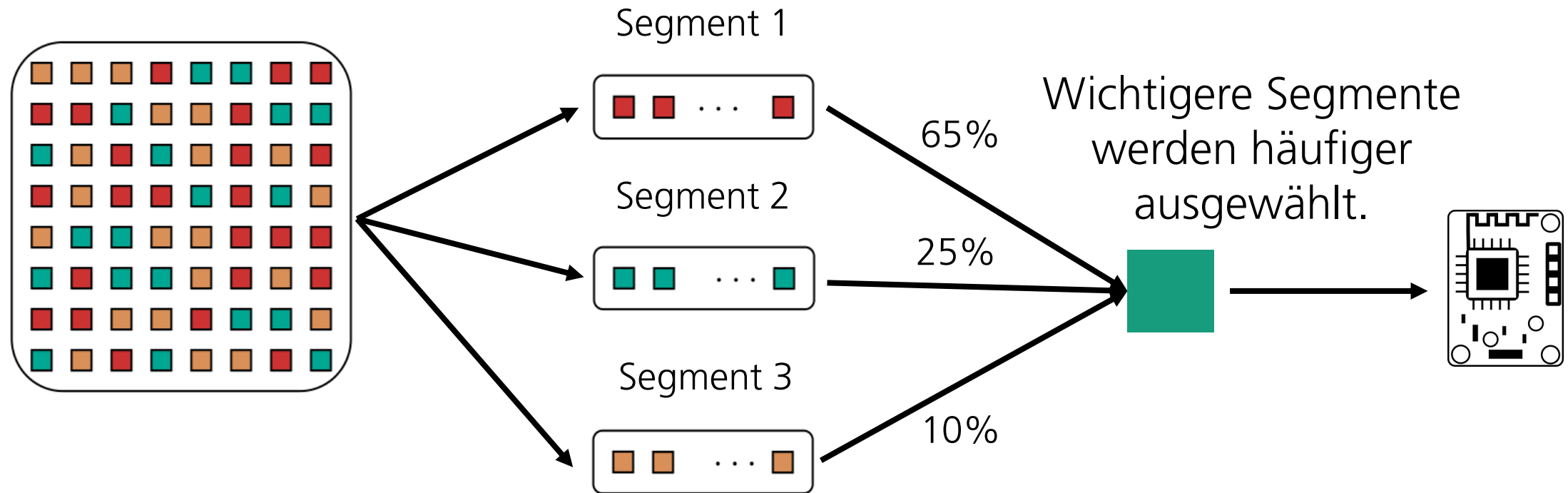
Probabilistisches Teilen



Bessere Vielfalt, jedoch wirkt es wie zufällige Segmentierung oder sogar noch unvorteilhafter!

Gist

Probabilistisches Teilen

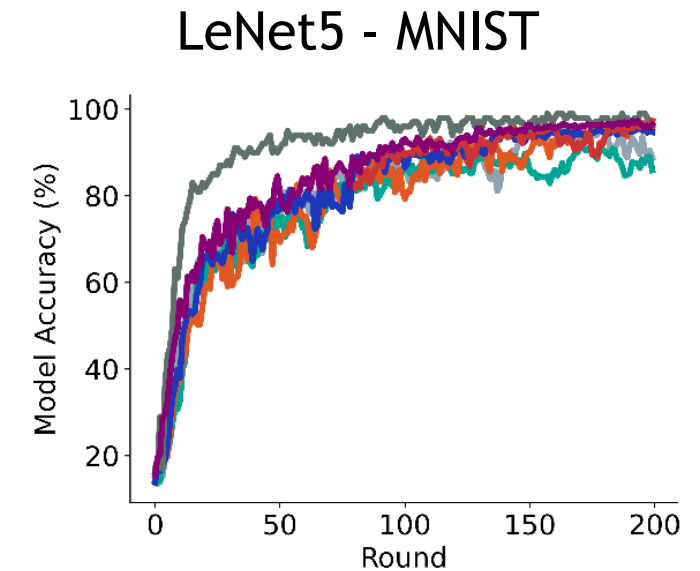
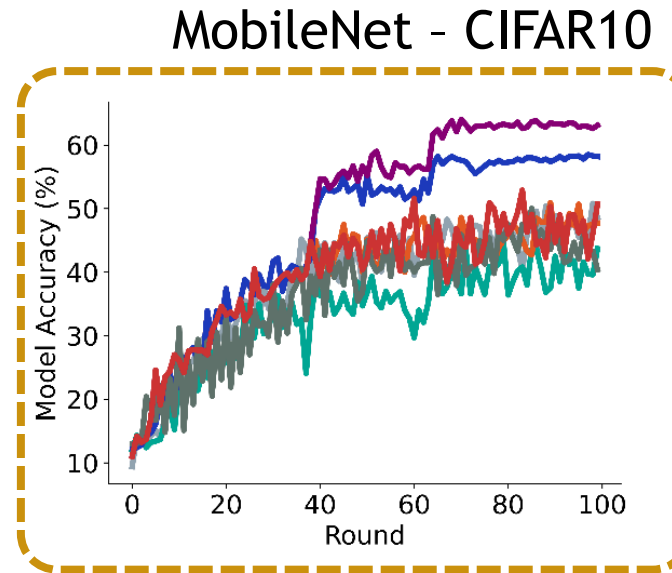
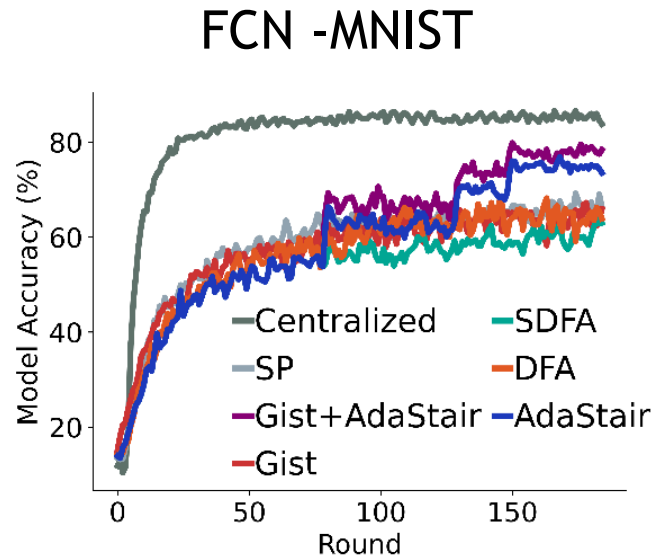


Wichtige Segmente verwenden, weniger wichtige trotzdem erkunden.



Gist

Wie schneidet **Gist** bei verschiedenen Modellen und Datensätzen ab?



Gist > Basislinien um ~4-5 % im Laufe der Zeit und bei Konvergenz.

Gist erzielt eine Steigerung von bis zu 20 %, wenn es mit Adastair kombiniert wird.

Gist hat einen größeren Einfluss auf komplexere Modelle.

Zusammenfassung

AIfES

Open Source ML Framework für
eingebettete Systeme



GitHub



Paper

Zusammenfassung

AiFES

Open Source ML Framework für eingebettete Systeme



GitHub



Paper

Gist

Ein dezentralisierter Segmentierungs- und Aggregationsansatz für Tiny Federated Learning auf Mikrocontrollern

- Höhere Genauigkeit und schnellere Konvergenz, insbesondere bei komplexeren Modellen.
- Weitere Details und Ergebnisse finden Sie im Paper.



Gist



Zwe³

Vielen Dank!

—

Dr. Lars Wulfert

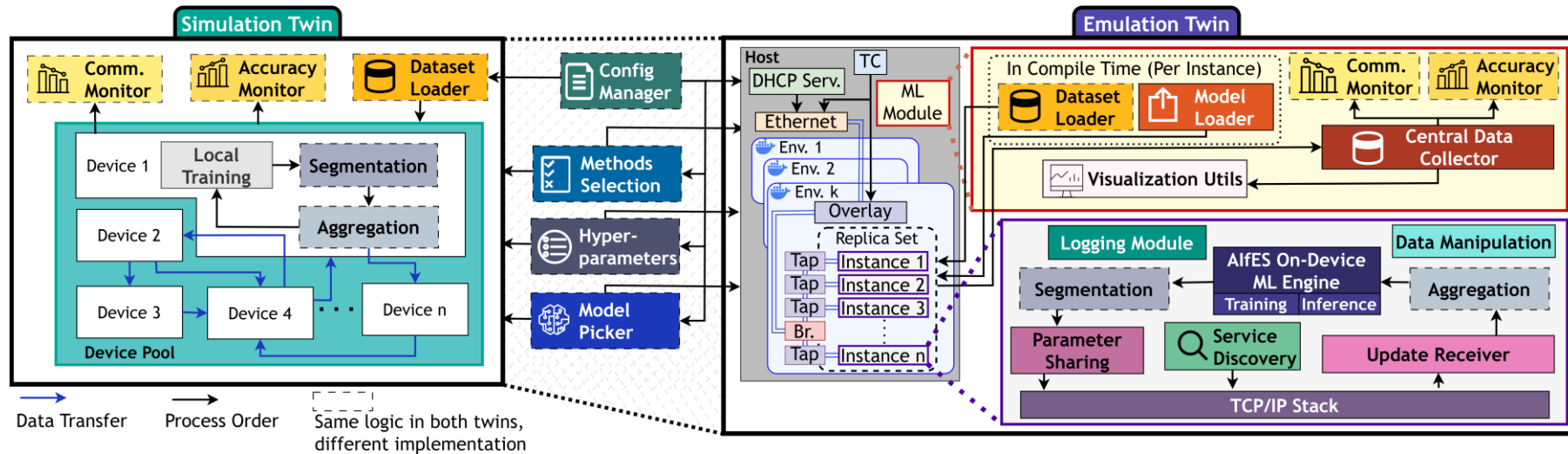
Smart Embedded Systems

lars.wulfert@ims.fraunhofer.de

www.aifes.ai

[Follow us on LinkedIn](#)





Same modules are deployed on real devices for measurement

Simulation helps faster iterations while emulation provides higher fidelity!

Asadi, Wulfert et al. "Road to Tiny Reality..." MobiCom Posters 2025

