

Unrolling the Performance of ZK-Rollups through Stochastic Modeling

Carlos Melo, **Johnnatan Messias**, José Miquéias, Glauber Gonçalves, Francisco A. Silva, and André Soares



IEEE SMC 2025, Vienna, Austria



MAX PLANCK INSTITUTE
FOR SOFTWARE SYSTEMS

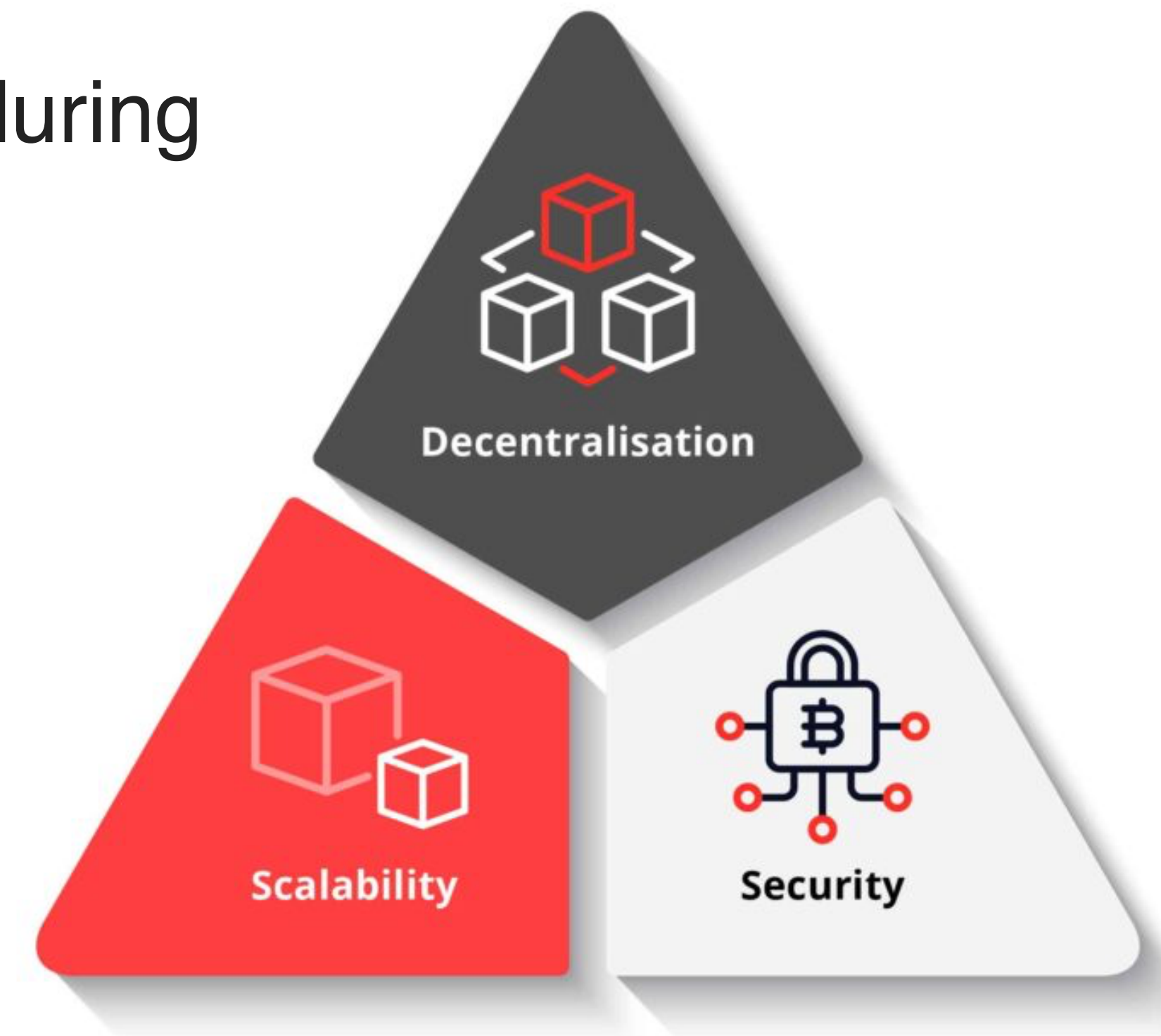
October 7th, 2025

johnnatan-messias.github.io

The Scalability Trilemma 🤝



- Blockchain technology offers data security and integrity.
- However, public blockchains like Ethereum face the “**Scalability Trilemma**”:
 - 😞 It’s hard to simultaneously optimize **Security**, **Decentralization** and **Scalability**.
- This results in low transaction throughput and high fees during periods of high demand, hindering broader adoption.



Why ZK-Rollups Matter



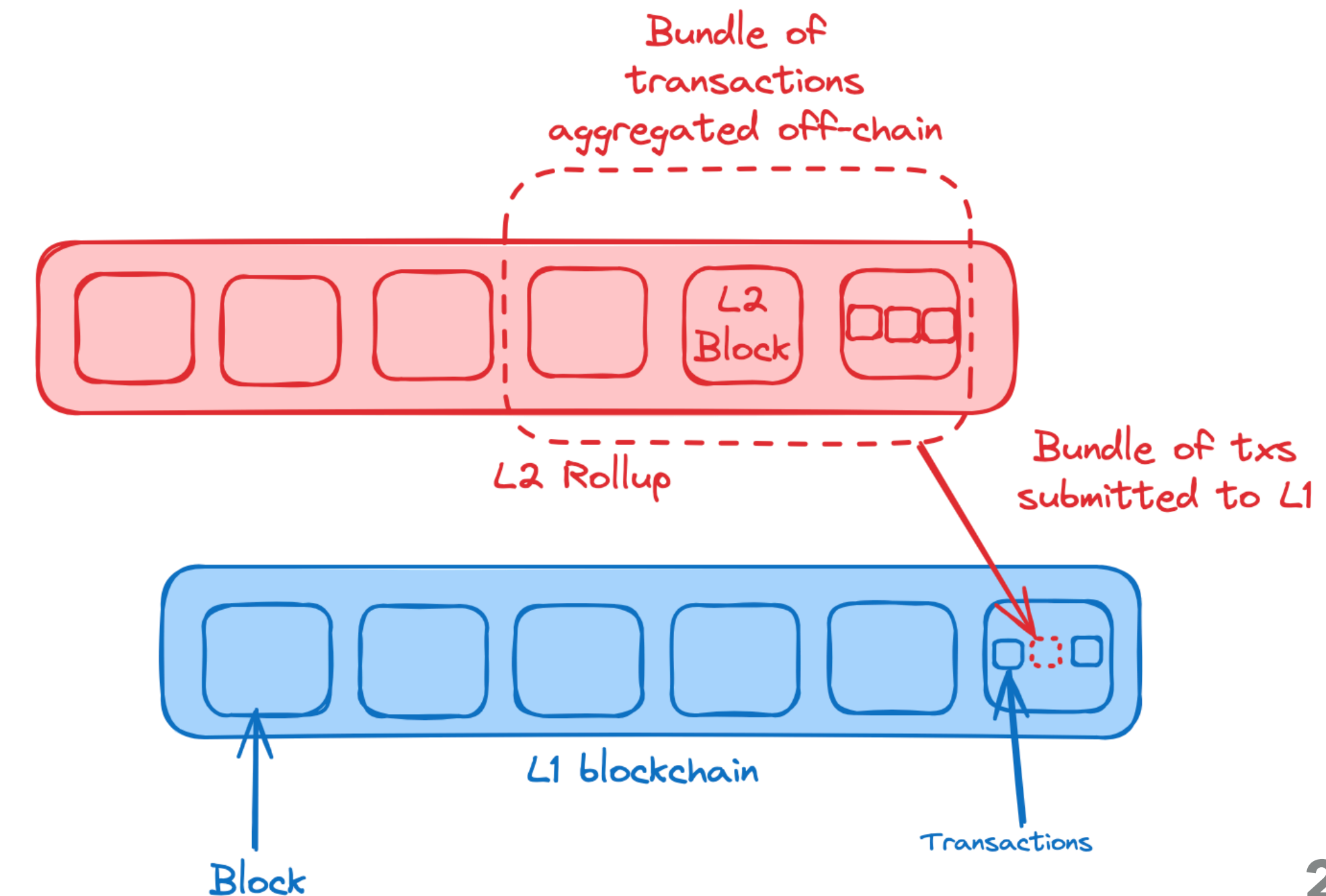
- **Layer-2** solutions emerged to solve this scalability challenge.

- **Rollups**

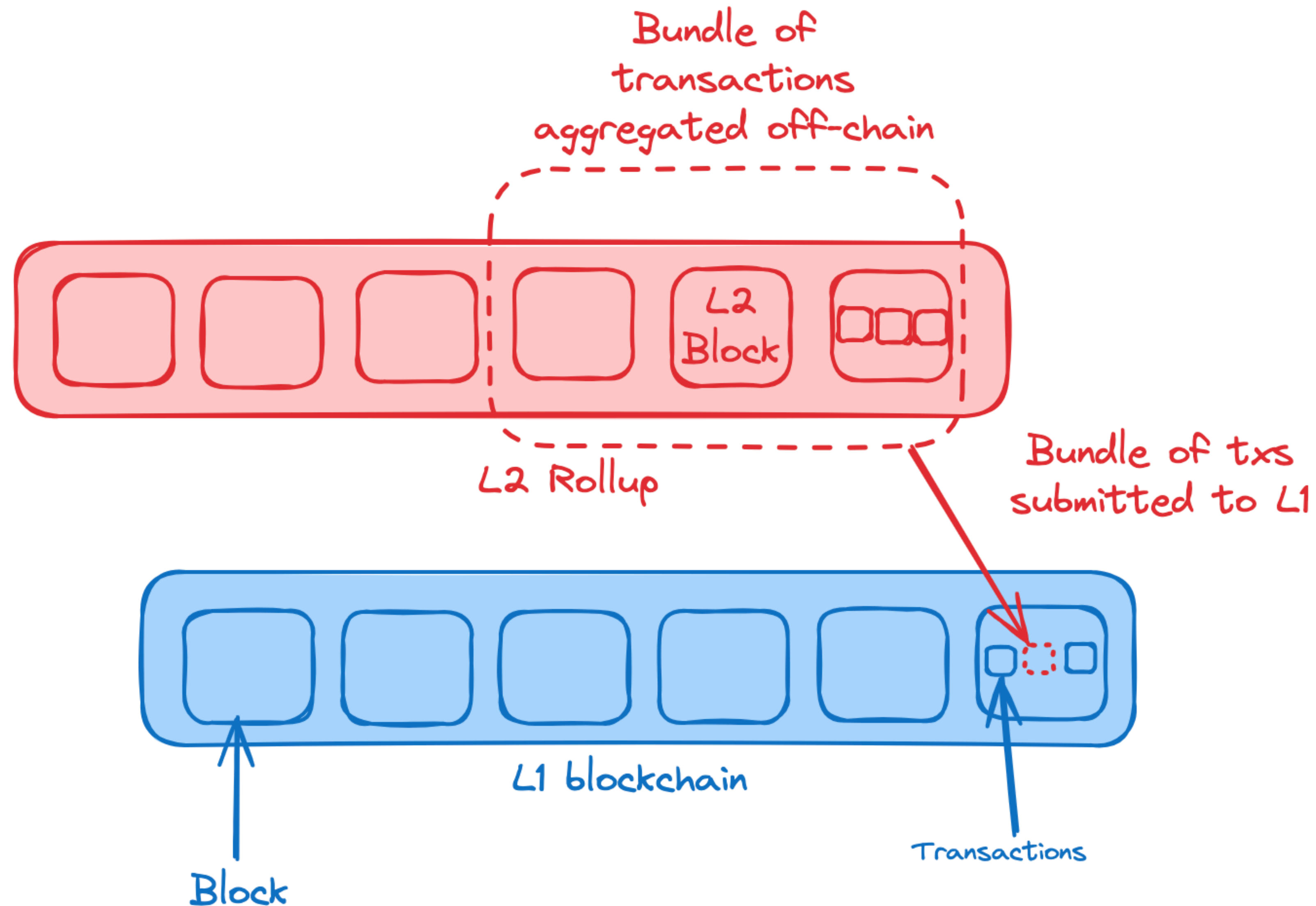
They process hundreds of transactions “off-chain”.

They post only essential data and a cryptographic proof “on-chain” to **Layer-1**.

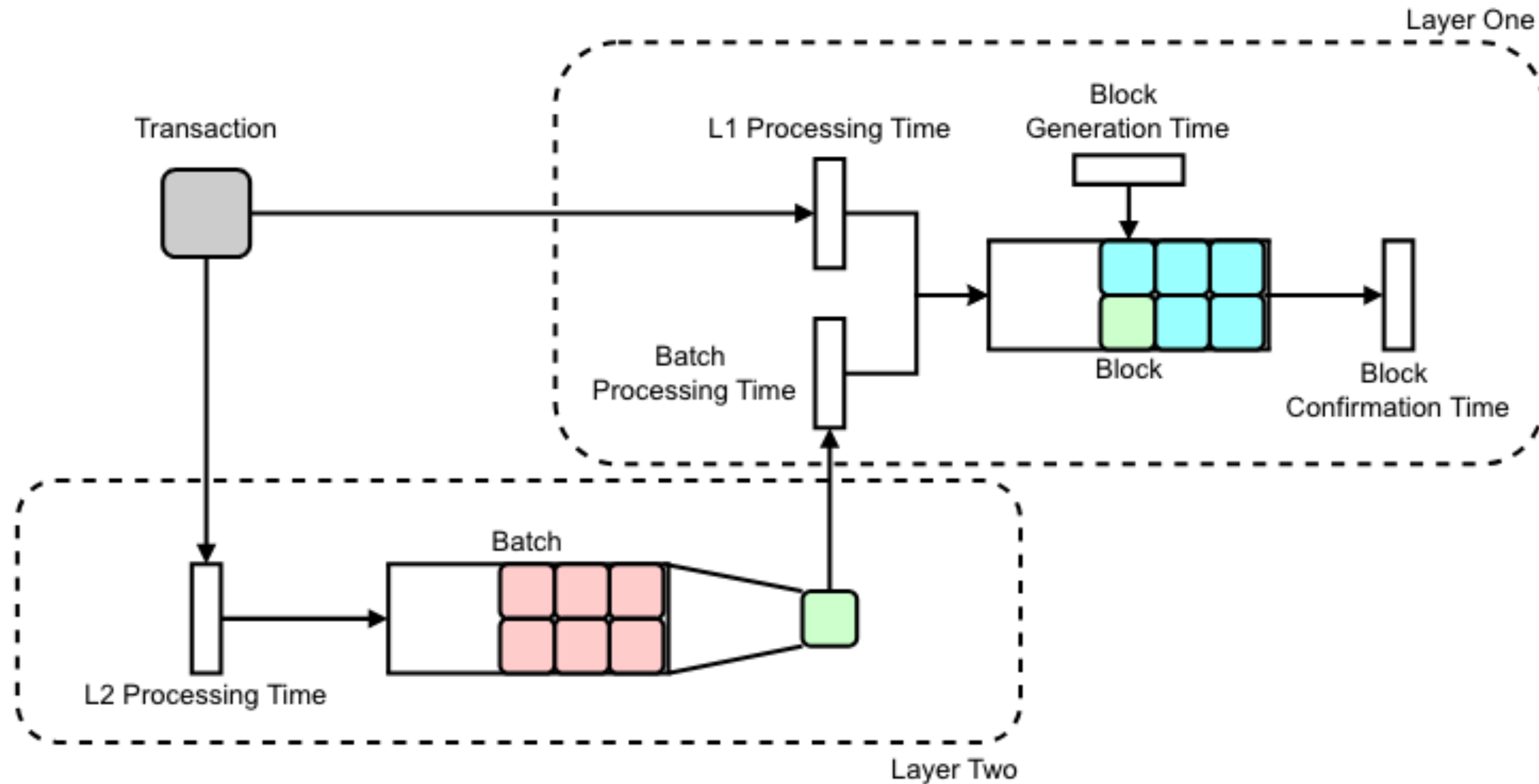
- **ZK-Rollups** use “zero-knowledge proofs” for validation, enhancing security and scalability.



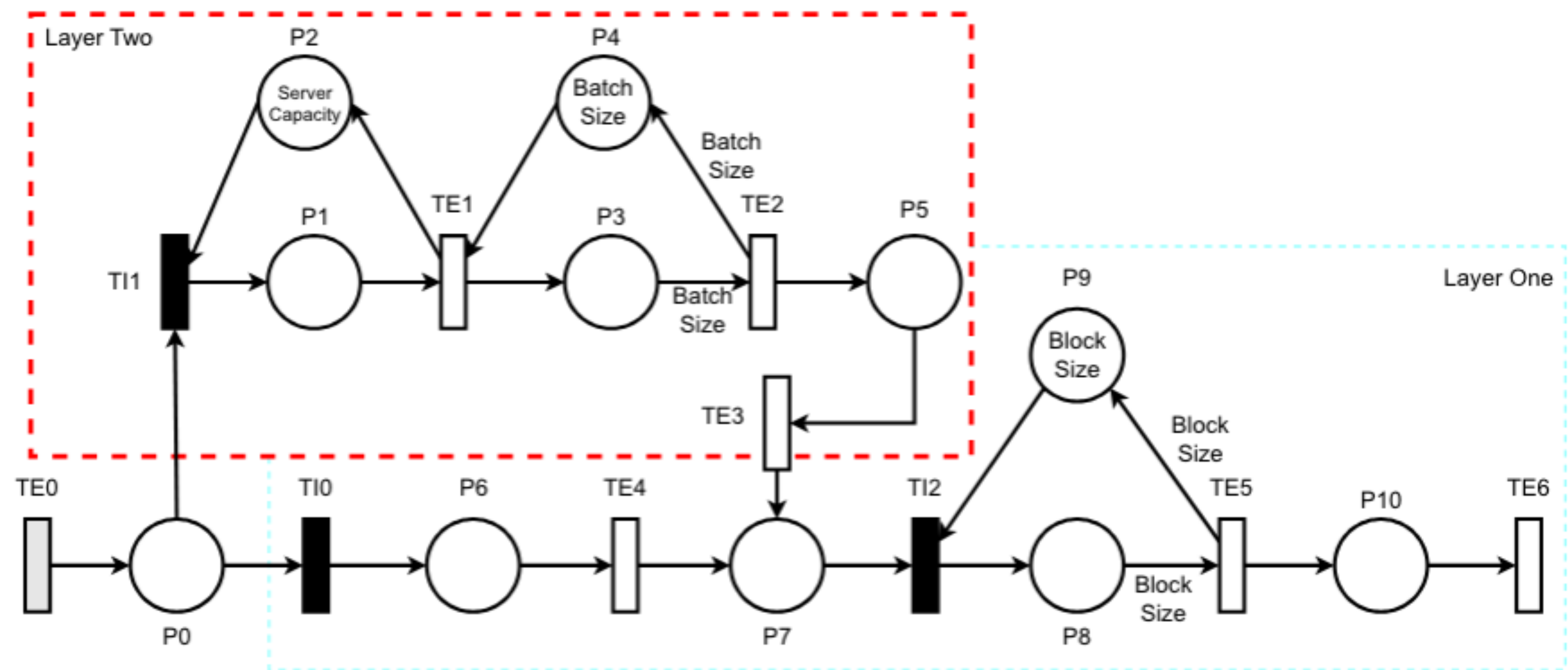
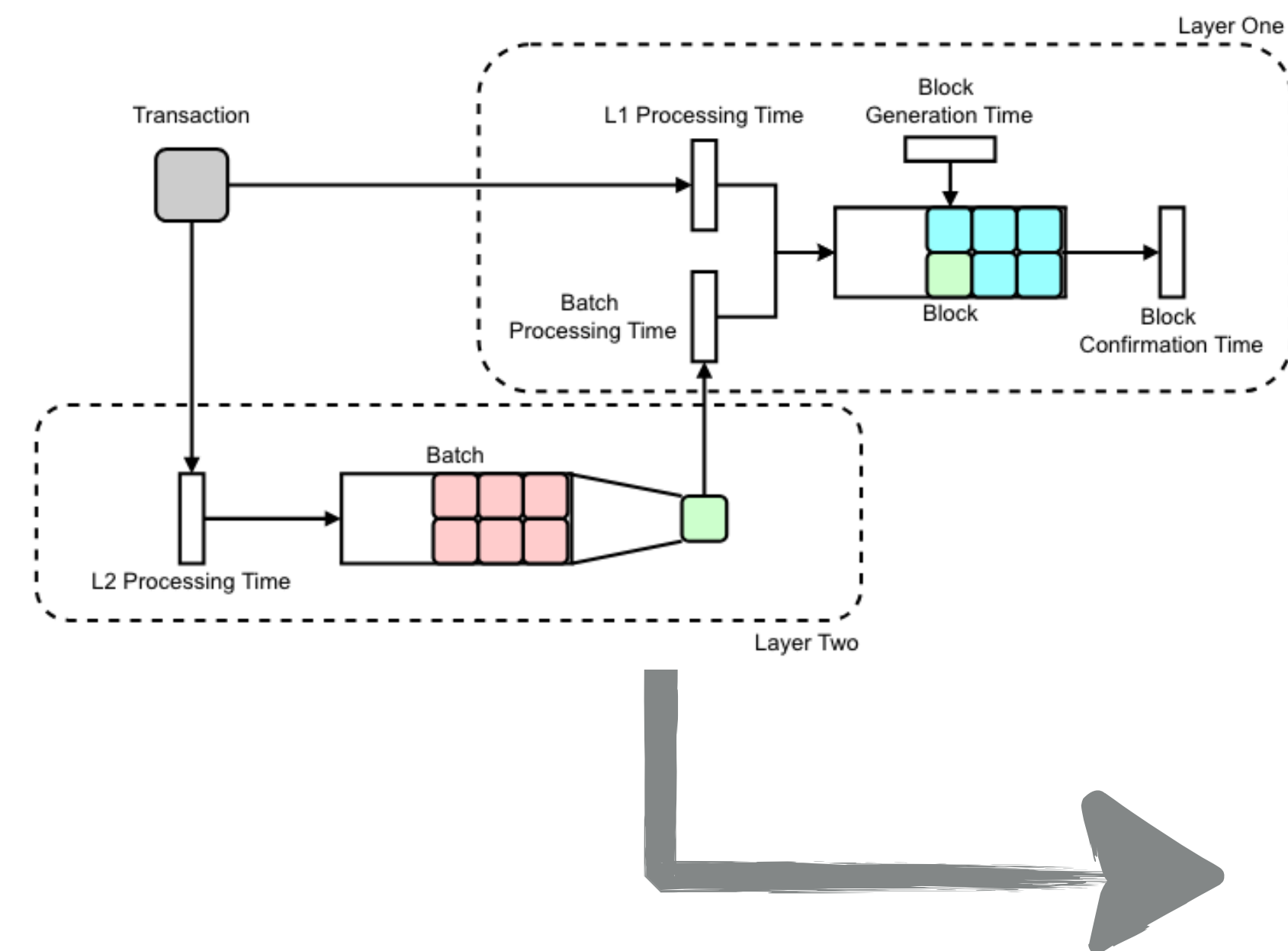
How Transactions Flow in Rollups



How Transactions Flow in Rollups



How Transactions Flow in Rollups



Our Case Studies



- **Case Study 1: System-Level** ⚡

Goal: Evaluate latency and throughput of the whole system as transactions are distributed between Layer-1 and Layer-2.

- **Case Study 2: Parameter Impact** ⚙️

Goal: Identify which Layer-2 parameters most affect performance.

- **Case Study 3: L2 Internals** 📦

Goal: Examine how batch processing time and batch size affect Layer-2 throughput directly.

Model Parameters & Inputs



Factor	Baseline	Variation {min; max}
Arrival (TE0)	~13.7 tps	{-}
L2 Probability	90 %	{0%; 100%}
L1 Probability	10 %	{0%; 100%}
Layer 1 Processing Time (TE 4)	~12.98 s	{-}
Layer 2 Processing Time (TE 1)	100 ms	{50 ms; 200 ms}
Block Generation Time (TE 5)	13 s	{-}
Block Confirmation Time (TE 6)	60 s	{-}
Batch Processing Time (TE 2)	1075 s	{600 s; 1800 s}
Batch Generation Time (TE 3)	1 s	{0.5 s; 1.5 s}
Batch Size	5000	{1000; 10,000}
Server Capacity	32	{16; 64}
Block Size	167	{-}

Input Parameters for
the Proposed Model

Model Parameters & Inputs

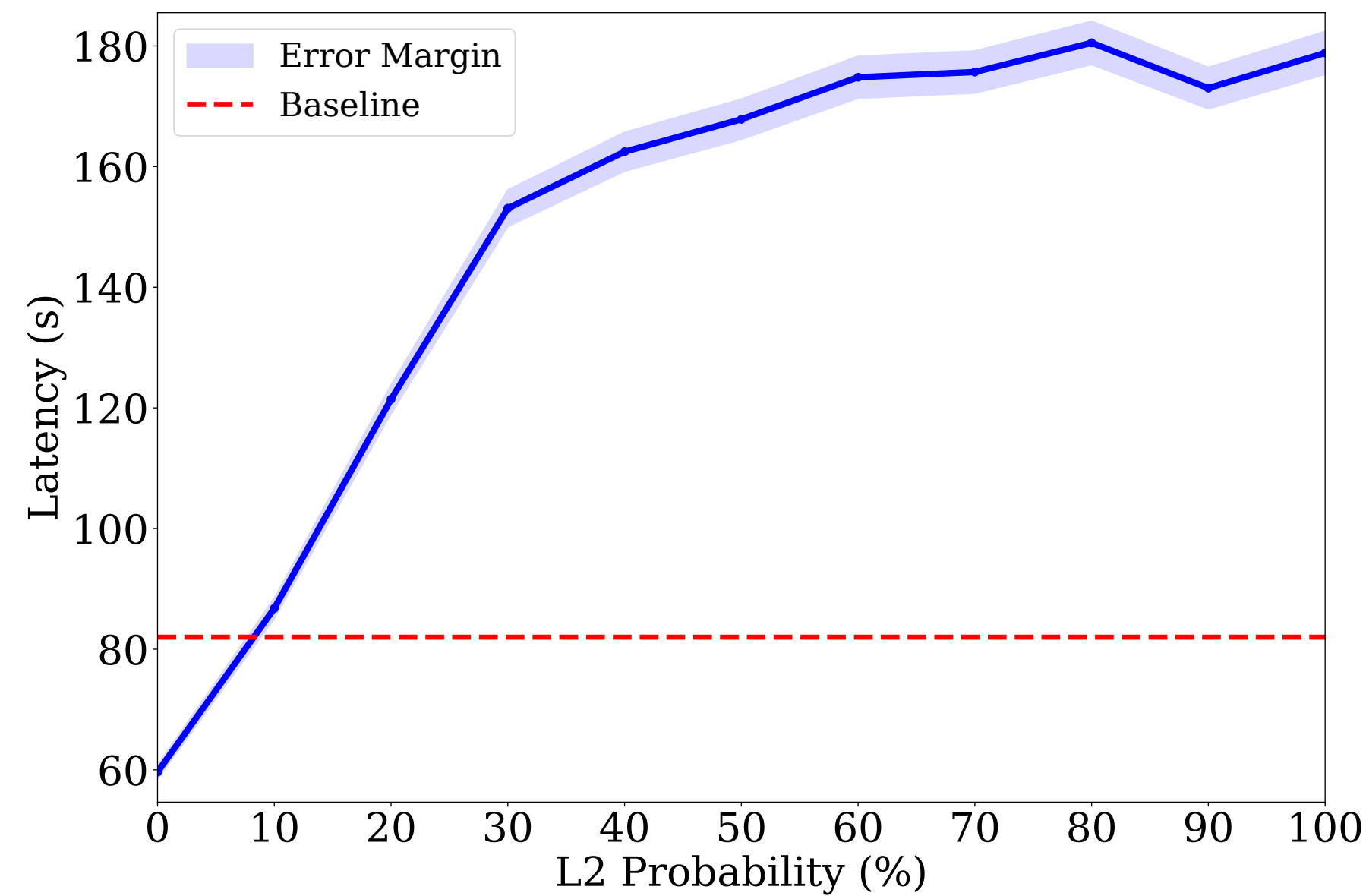
Obtained from real-world data!



Factor	Baseline	Variation {min; max}
Arrival (TE0)	~13.7 tps	{-}
L2 Probability	90 %	{0%; 100%}
L1 Probability	10 %	{0%; 100%}
Layer 1 Processing Time (TE 4)	~12.98 s	{-}
Layer 2 Processing Time (TE 1)	100 ms	{50 ms; 200 ms}
Block Generation Time (TE 5)	13 s	{-}
Block Confirmation Time (TE 6)	60 s	{-}
Batch Processing Time (TE 2)	1075 s	{600 s; 1800 s}
Batch Generation Time (TE 3)	1 s	{0.5 s; 1.5 s}
Batch Size	5000	{1000; 10,000}
Server Capacity	32	{16; 64}
Block Size	167	{-}

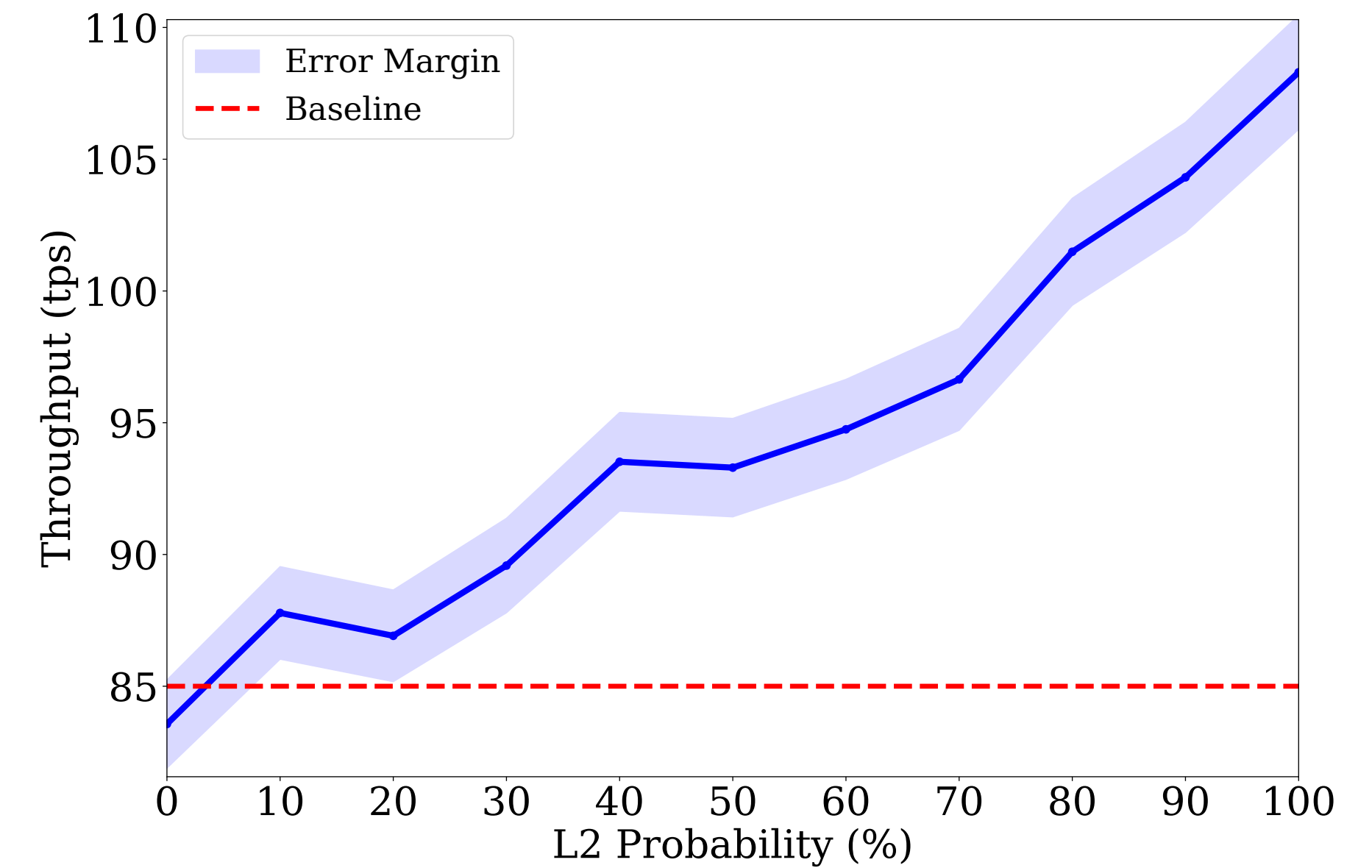
Input Parameters for the Proposed Model

Case Study 1: Throughput vs Latency Trade-Off



Negative impact of Layer-2 on system Latency

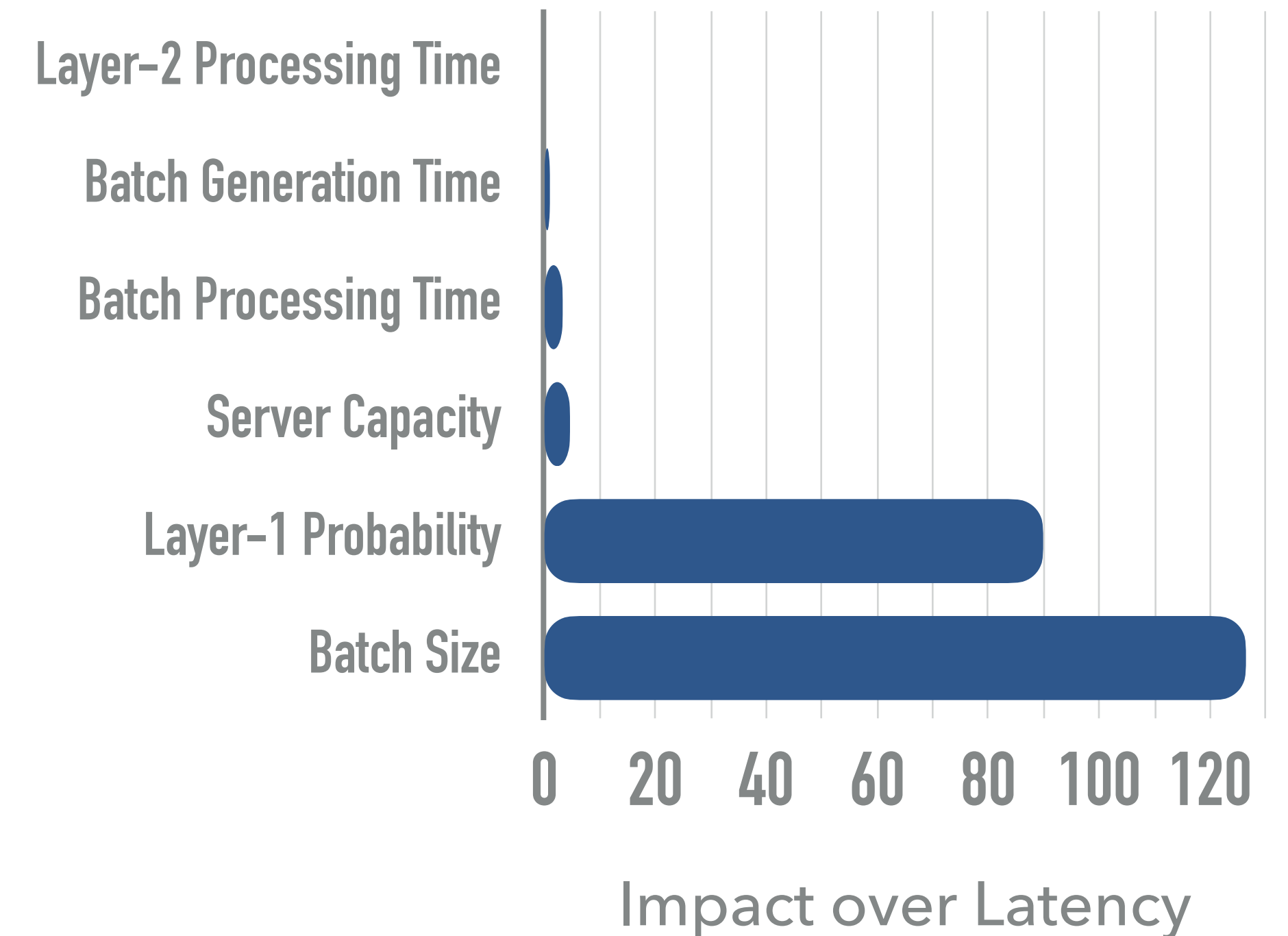
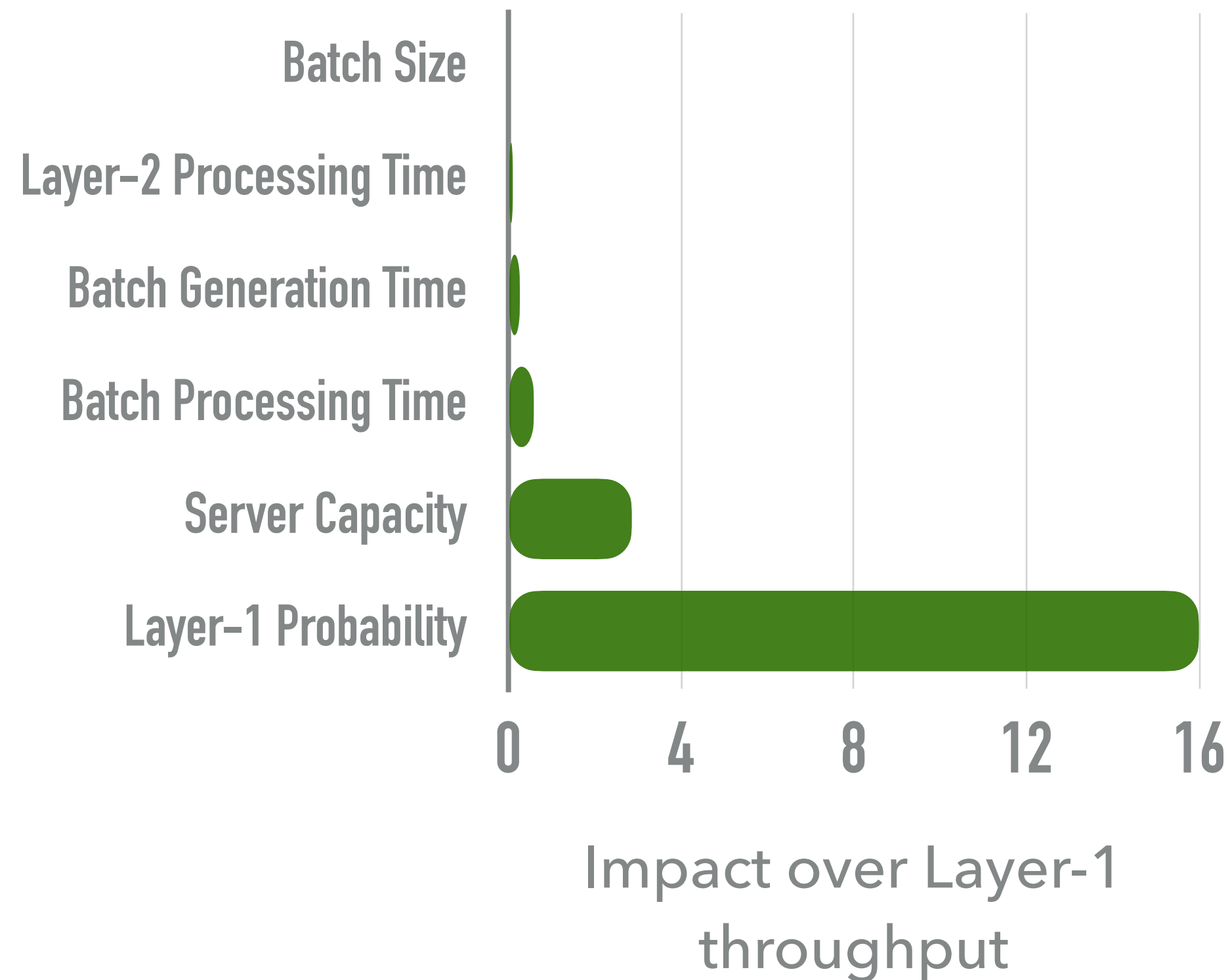
Latency and Throughput vs. Probability of transactions using L2



Positive impact of Layer-2 on general throughput

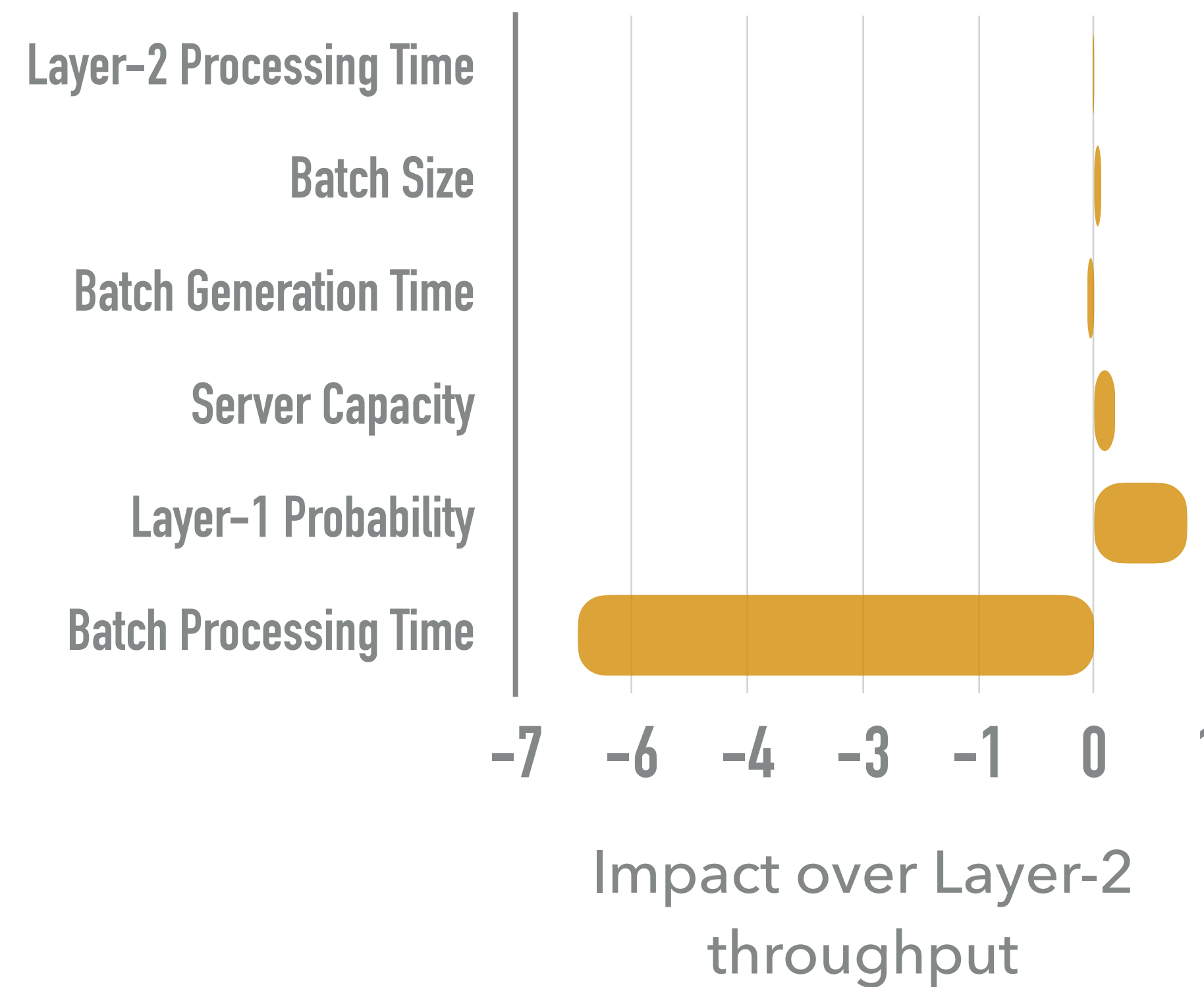
- ➡ More transactions routed to Layer-2 → higher throughput (up to ~20% improvement).
- ➡ But, more Layer-2 transactions → much higher latency (over +100%).
- ➡ **Implication:** Users face a trade-off between faster global throughput and slower transaction confirmation.

Case Study 2: What Drives Performance?



- **Throughput:** Strongly affected by the probability of using Layer-1.
- **Latency:** Strongly impacted by batch size (adds 120+ seconds).
- **Implication:** Batch size is the dominant driver of latency, while Layer-1 vs Layer-2 distribution drives throughput.

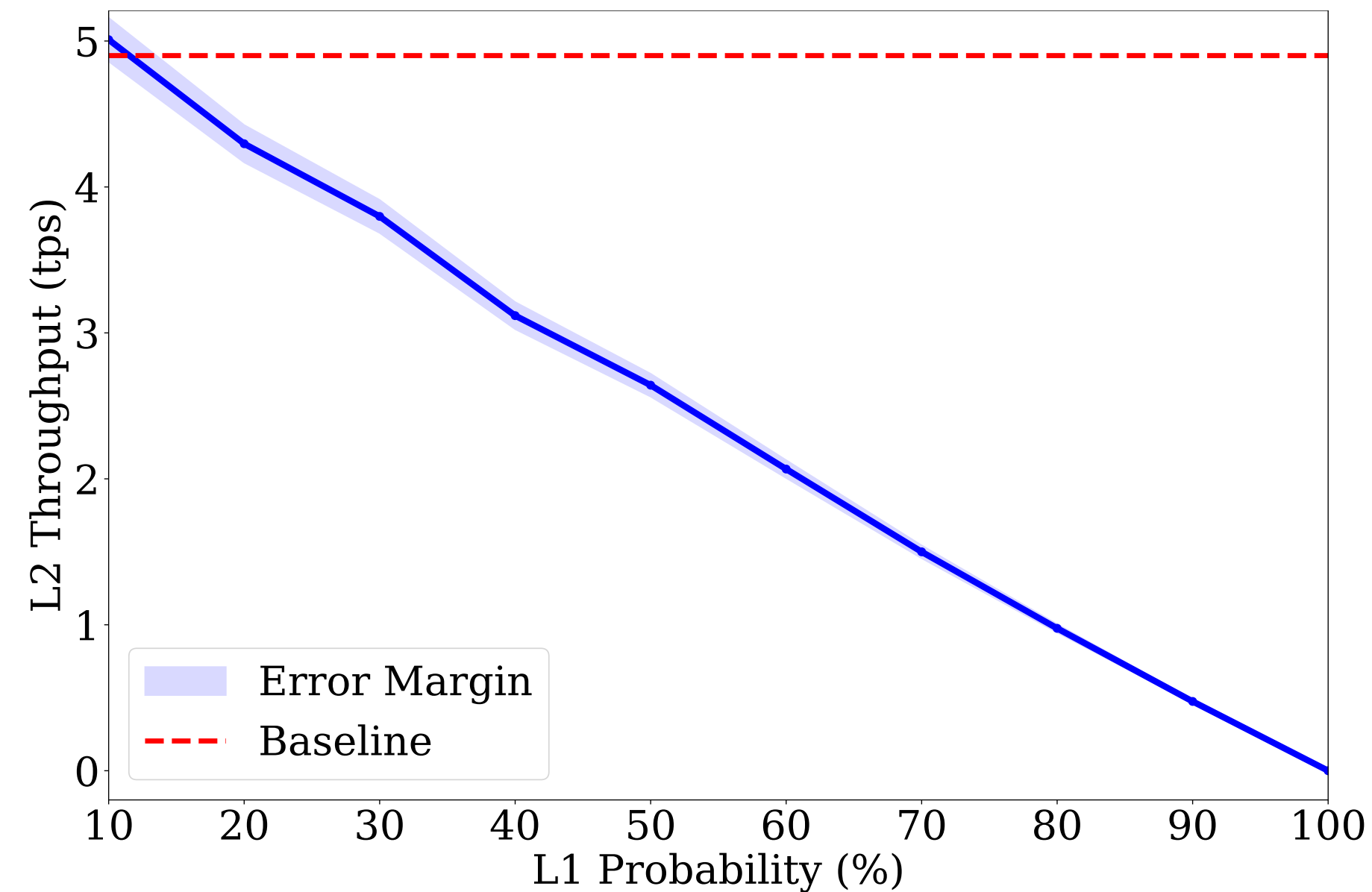
Case Study 2 (Cont.): Bottlenecks Inside Layer-2



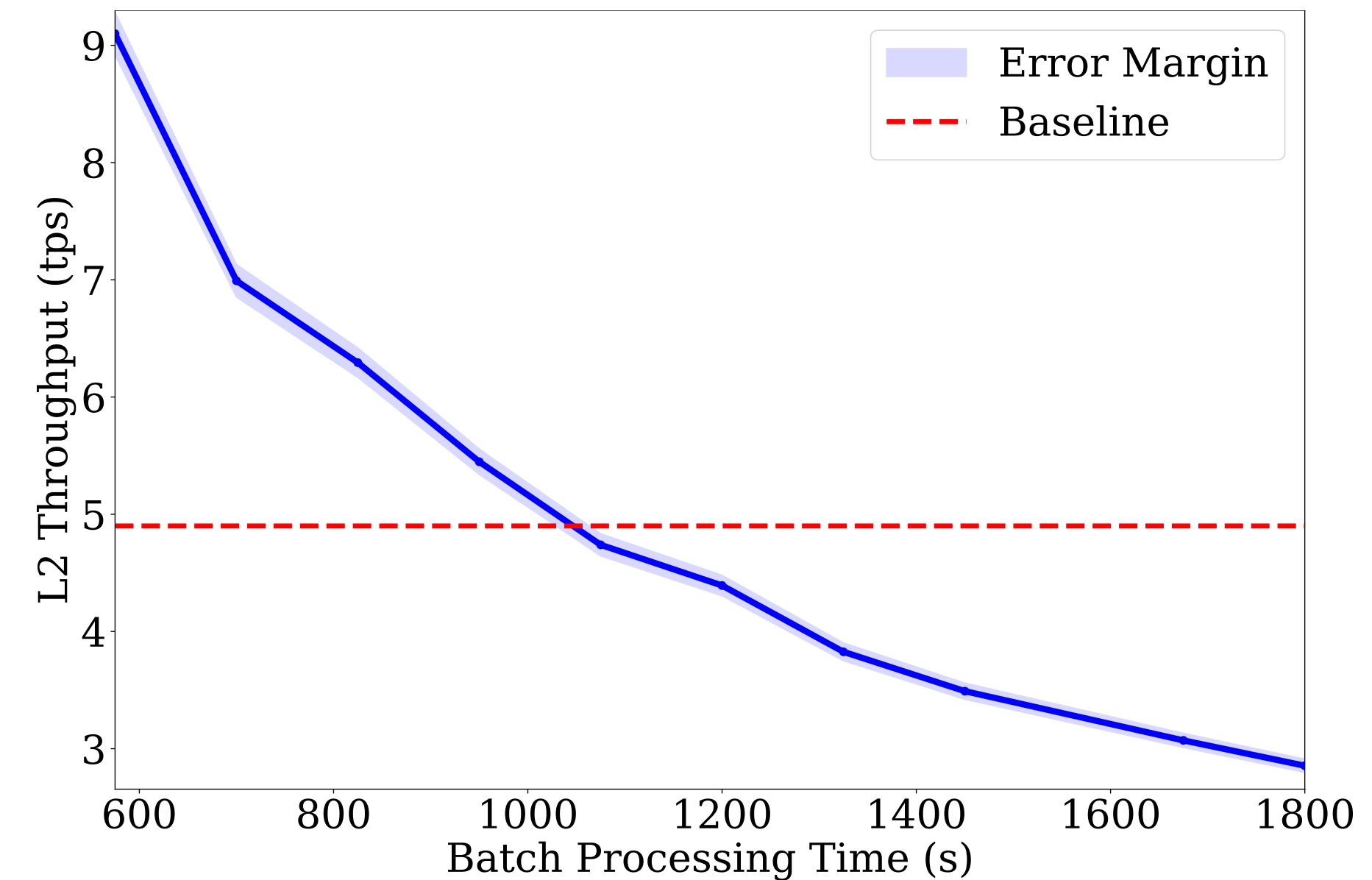
High impact of batch processing time on Layer-2

- **Batch processing** time strongly reduces throughput (> 6 tps lost).

Case Study 3: Batch Processing vs Probability



Layer-2 Throughput vs. Probability of Layer-2 and Batch Processing Time



- L2 throughput increases as more transactions use Layer-2.
- Batch processing time is critical.
- **Implication:** Smaller batches improve Layer-2 throughput but at higher costs (more frequent proof submissions).

Key Takeaways & Future Work



● Contribution

- Introduced a **Stochastic Petri Net model** to analyze Layer-1 & Layer-2 interactions.

● Findings

- ZK-Rollups boost throughput (+20%) while keeping security & decentralization.
- ⚖️ Larger batches → latency > 2×, showing a fundamental trade-off.

● Future work

- ✅ Validate with real ZK-Rollup deployments.
- 🔄 Extend to Optimistic Rollups.
- 📦 Explore adaptive batch sizing to balance cost, latency, and throughput.

Let's Connect

johnme@mpi-sws.org
johnnatan-messias.github.io



Johnnatan Messias, PhD
Research Scientist

  @johnnatan_me



MAX PLANCK INSTITUTE
FOR SOFTWARE SYSTEMS