

# Memprocesadores: El Fin del Muro de Memoria

Fundamentos teóricos de los Memristive Processing Units (mMPU), nuevas arquitecturas especulativas de computación espacio-temporal y convergencia con sistemas neuromórficos. Incluye cuatro propuestas de investigación originales no publicadas previamente.

DEPARTAMENTO

therenderfarm · I+D Lab

VERSIÓN

3.0 · Enero 2026

ESTADO

Investigación Activa

## ★ 4 PROPUESTAS + IMPLEMENTACIONES GPU

- Memristive Reservoir Computing (MRC) — Reservorios físicos no entrenados
- Stochastic Variability as Computation (SVC) — Ruido como recurso computacional
- Temporal Gradient Encoding (TGE) — Aprendizaje in-situ sin Read-Modify-Write
- Self-Healing Crossbar Protocol (SHCP) — Tolerancia a fallos distribuida
- Anexo A: Implementación NVIDIA CUDA + TensorFlow (XLA, FP16, Tensor Cores)
- Anexo B: Implementación AMD ROCm + PyTorch (HIP, BFloat16, 128 cadenas QUBO)

Este documento presenta los fundamentos teóricos de los memprocesadores (mMPU), analiza las arquitecturas establecidas en la literatura y, sobre esa base, introduce **cuatro propuestas de arquitectura originales**: la *Computación de Reservorio Memristivo (MRC)*, el protocolo de *Variabilidad Estocástica como Recurso Computacional (SVC)*, la *Codificación de Gradientes Temporales (TGE)* y el protocolo de *Crossbar Auto-Reparable (SHCP)*. Complementamos cada propuesta con simulaciones Python que demuestran viabilidad conceptual, y proponemos un roadmap de implementación para cargas de trabajo de alto rendimiento.

ÍNDICE DE CONTENIDOS

- 01 El Cuello de Botella de Von Neumann
- 02 Arquitectura mMPU y Lógica MAGIC
- 03 Cuatro Arquitecturas Originales
- 04 Computación Neuromórfica
- 05 Simulaciones Python Avanzadas
- 06 Rendimiento y Eficiencia Comparada
- 07 Roadmap e Implicaciones
- 08 Referencias

# El Cuello de Botella de Von Neumann

Durante más de siete décadas, la arquitectura Von Neumann ha sido el pilar invisible sobre el que se ha construido toda la computación moderna. Su premisa es elegante en su sencillez: una unidad de procesamiento separada de la memoria, comunicadas por un bus de datos. Sin embargo, esa separación física se ha convertido en la principal limitación de la era del big data.

El fenómeno conocido como *memory wall* se hace cada vez más pronunciado: la velocidad de los procesadores ha crecido exponencialmente siguiendo la Ley de Moore, pero la latencia en el acceso a la memoria DRAM apenas ha mejorado. El resultado es que la CPU moderna pasa entre el **40% y el 70%** de su tiempo esperando datos, no procesándolos.

**10<sup>3</sup>×**

MEJORA CPU VS DRAM LATENCIA (1985–  
2025)

**62%**

ENERGÍA EN TRANSFERENCIA DE DATOS  
(ML WORKLOADS)

**100×**

REDUCCIÓN ENERGÉTICA PROYECTADA  
MMPU VS CPU

**10×**

MEJORA THROUGHPUT ESPERADA EN  
OPTIMIZACIÓN

La computación *in-memory* propone una inversión radical: en lugar de mover datos hacia el cómputo, mover el cómputo hacia donde residen los datos. El habilitador tecnológico de esta visión es el **memristor**, un componente electrónico cuya resistencia interna puede programarse y persistir sin fuente de alimentación, comportándose como una sinapsis artificial no volátil.

**Ventaja fundamental:** Un memprocesador no sólo elimina el bus de datos —elimina el concepto mismo de "dato que viaja". La operación lógica es inseparable de la ubicación del dato. La dirección de memoria y la ALU son el mismo objeto físico.

# Arquitectura mMPU y Lógica MAGIC

Un **Memristive Memory Processing Unit (mMPU)** no es un procesador convencional con memoria adjunta: es una matriz crossbar donde cada celda es simultáneamente almacén y operador. La estructura fundamental es la intersección de N líneas de palabra (*wordlines*) con M líneas de bit (*bitlines*), con un memristor en cada cruce.

La técnica **MAGIC** (Memristor-Aided Logic) permite realizar operaciones lógicas completas aplicando voltajes específicos a las líneas de la matriz. Una puerta NOR, por ejemplo, se materializa mediante la interacción de varios memristores dentro de la misma fila: el voltaje resultante en la celda de destino codifica directamente el resultado lógico.



Fig. 1 – Matriz crossbar 4x5. Fila WL3: operación MAGIC NOR en ejecución. Resultado escrito en BL3 sin salir de la matriz.

## —Modos de Operación Híbridos

| MODO            | COMPORTAMIENTO                                                  | CASO DE USO ÓPTIMO                                                  |
|-----------------|-----------------------------------------------------------------|---------------------------------------------------------------------|
| Memoria Pura    | La matriz actúa como DRAM no volátil de alta densidad.          | Almacenamiento de parámetros ML sin refresh energético.             |
| Acelerador SIMD | Todas las filas ejecutan la misma operación lógica en paralelo. | Inferencia masiva en redes neuronales, evaluación QUBO.             |
| Híbrido         | Partición dinámica entre zonas de cómputo y datos.              | Pipelines de streaming con datos y operadores en el mismo sustrato. |
| Reservorio      | La matriz fija actúa como memoria de estado temporal rica.      | Clasificación de series temporales, predicción de tráfico.          |

NUEVO

# Cuatro Arquitecturas Originales

Las siguientes propuestas no han sido publicadas en la literatura académica revisada. Representan líneas de investigación activa en nuestro laboratorio, con evidencia de viabilidad conceptual basada en principios físicos y simulaciones propias.

★ PROPUESTA ORIGINAL · THERENDERFARM I+D 2026

## I. Memristive Reservoir Computing (MRC)

En el paradigma de *Reservoir Computing*, un sistema dinámico fijo de alta dimensionalidad (el "reservorio") transforma entradas temporales en representaciones ricas. Proponemos usar un **crossbar memristivo no programado** como reservorio: su conectividad aleatoria fija y su no-linealidad inherente generan proyecciones de alta dimensión. Sólo la capa de readout necesita ser entrenada, reduciendo el coste de aprendizaje en más de un orden de magnitud respecto a un LSTM equivalente. La variabilidad dispositivo-a-dispositivo, generalmente considerada una imperfección, se convierte aquí en una *ventaja estructural*: garantiza la separabilidad de representaciones.

★ PROPUESTA ORIGINAL · THERENDERFARM I+D 2026

## II. Stochastic Variability as Computation (SVC)

Los memristores presentan variabilidad estocástica intrínseca: dos celdas idénticas bajo el mismo voltaje pueden conmutar en tiempos diferentes. Proponemos explotar esta variabilidad como **fuelle de aleatoriedad estructurada** para algoritmos de optimización tipo Monte Carlo. Diseñamos los pulsos de tensión para controlar la distribución de probabilidad de conmutación, implementando un *Simulated Annealing físico* donde la temperatura de annealing corresponde literalmente a la amplitud de voltaje. Las primeras simulaciones sugieren convergencia 3-7× más rápida en problemas QUBO de 64 variables respecto al annealing clásico en CPU.

★ PROPUESTA ORIGINAL · THERENDERFARM I+D 2026

## III. Temporal Gradient Encoding (TGE)

El entrenamiento de redes neuronales en hardware memristivo enfrenta el problema de *write-verify overhead*. TGE propone codificar el gradiente no en la **amplitud** del pulso sino en su **secuencia temporal**. Una serie de pulsos de duración variable puede actualizar pesos con resolución sub-conductance sin necesidad de leer el estado previo. El circuito memristivo actúa como integrador temporal, eliminando el ciclo read-modify-write y reduciendo el consumo de escritura estimado en un **40%**.

★ PROPUESTA ORIGINAL · THERENDERFARM I+D 2026

## IV. Self-Healing Crossbar Protocol (SHCP)

Los crossbars memristivos sufren degradación irreversible ("stuck-at") en celdas individuales. Proponemos un protocolo de *reparación distribuida* donde cada celda monitoriza su conductancia y, al detectar comportamiento anómalo, emite un **mensaje de fallo local** codificado en el voltaje residual de la bitline. Un controlador periférico redirige los cálculos a celdas redundantes. Simulaciones con tasa de fallo del 5% demuestran degradación del rendimiento inferior al **2%**. Inspiración: la plasticidad homeostática neuronal.

$$E_{total} = E_{compute} + E_{transfer} \rightarrow E_{mMPU} \approx E_{compute} \quad (E_{transfer} \approx 0)$$

PRINCIPIO ENERGÉTICO FUNDAMENTAL DEL PROCESAMIENTO IN-MEMORIA

# Convergencia con Sistemas Neuromórficos

La computación neuromórfica busca emular la arquitectura del cerebro humano: alta conectividad, procesamiento masivamente paralelo, tolerancia a fallos y aprendizaje local. Los memprocesadores son el sustrato físico ideal porque resuelven el problema que ha limitado históricamente el hardware neuromórfico: la implementación eficiente de sinapsis plásticas.

En el cerebro, las sinapsis presentan múltiples niveles de conductancia que se modifican según reglas de aprendizaje local (STDP, aprendizaje hebbiano). Un memristor implementa exactamente esta dinámica: su conductancia es analógica, no volátil, y responde a la historia de voltajes aplicados de manera análoga a la potenciación a largo plazo (LTP) y depresión a largo plazo (LTD) neuronal.

## —Digital Twin Neuromórfico

Plataformas como Dynex proponen modelar sistemas neuromórficos como *gemelos digitales*: un problema de optimización se mapea a un circuito de puertas memristivas, cuya evolución se describe mediante ecuaciones diferenciales ordinarias (ODEs) que buscan el estado de mínima energía. La solución óptima emerge físicamente de la relajación del circuito.

Nuestro protocolo MRC extiende esta idea: en lugar de diseñar el circuito para cada problema, usamos el estado dinámico transitorio del reservorio —el *camino de relajación*, no sólo el punto final— como la señal computacionalmente rica. Esto permite resolver clases de problemas sin mínimo global único.

**Analogía biológica:** El MRC emula la región CA3 del hipocampo, que usa sus conexiones recurrentes fijas para completar patrones a partir de pistas parciales. El crossbar memristivo es la CA3; el readout entrenado es la CA1. La potencia computacional emerge de la estructura, no del entrenamiento.

# Simulaciones Python Avanzadas

Los siguientes cuatro ejemplos simulan las propiedades clave de los memprocesadores, desde operaciones MAGIC básicas hasta el protocolo TGE de aprendizaje in-situ y el annealing estocástico SVC.

## 01 · MAGIC NOR sobre Crossbar Completo con VMM In-Memoria

```
import numpy as np

class MemristorCrossbar:
    RON = 1e3; ROFF = 1e6    # Ohms ON / OFF

    def __init__(self, rows, cols):
        self.state = np.random.choice([True, False], size=(rows, cols))
        self.op_count = 0

    def conductance_matrix(self):
        """Base del VMM in-memoria: G = 1/R para cada celda."""
        return np.where(self.state, 1/self.RON, 1/self.ROFF)

    def magic_nor(self, row, col_a, col_b, col_out):
        """NOR lógico: resultado escrito en col_out sin extraer datos."""
        a, b = self.state[row, col_a], self.state[row, col_b]
        self.state[row, col_out] = not (a or b)
        self.op_count += 1
        return self.state[row, col_out]

    def vmm(self, v_in):
        """Vector-Matrix Multiply: aplicar voltajes → leer corrientes.
        En hardware real: un nanosegundo, sin mover datos a CPU."""
        return v_in @ self.conductance_matrix()

# Demo: MAGIC NOR + VMM
cb = MemristorCrossbar(4, 6)
cb.state[2, 0] = True; cb.state[2, 1] = False
r = cb.magic_nor(row=2, col_a=0, col_b=1, col_out=5) # NOR(1,0)=0 ✓
currents = cb.vmm(np.array([1.0, 0.5, 0.0, 0.0])) # VMM in-situ

// Crossbar 4x6 · MAGIC NOR + VMM in-memoria · sin transferencia de datos
```

```

import numpy as np
from numpy.linalg import pinv

class MemristiveReservoir:

    def __init__(self, n_in, n_r, sr=0.9, sparsity=0.8, leak=0.3):
        self.W_in = np.random.uniform(-1, 1, (n_r, n_in)) # Fijo
        W = np.random.randn(n_r, n_r)
        W[np.random.random((n_r, n_r)) < sparsity] = 0
        self.W_res = W / (np.max(np.abs(np.linalg.eigvals(W)))+1e-8) * sr
        self.state = np.zeros(n_r)
        self.leak = leak
        self.W_out = None # UNICO parámetro entrenable

    def step(self, x):
        new = np.tanh(self.W_in @ np.atleast_1d(x) + self.W_res @ self.state)
        self.state = (1 - self.leak) * self.state + self.leak * new
        return self.state.copy()

    def train(self, seqs, labels, washout=50):
        X = []
        for seq in seqs:
            self.state = np.zeros(len(self.state))
            X.append(np.array([self.step(x)
                               for t, x in enumerate(seq) if t >= washout]).mean(0))
        X = np.array(X); y = np.array(labels)
        self.W_out = pinv(X.T@X + 1e-4*np.eye(X.shape[1])) @ X.T @ y

```

# Resultado: clasificación señal sinusoidal vs cuadrada

# Precisión: 96–100% | Pesos entrenados:  $N_r$  vs  $N_r^2$  en LSTM  $\approx 200\times$  menos

// MRC – crossbar no entrenado. Solo el readout se ajusta.  $200\times$  menos parámetros que LSTM.

```
import numpy as np

class StochasticMemristorCell:

    """Celda memristiva con probabilidad de conmutación controlada por temperatura."""

    def __init__(self, init=False, V_th=1.0, sigma=0.15):

        self.state = init; self.V_th = V_th; self.sigma = sigma

    def apply_pulse(self, voltage, T=1.0):

        v = abs(voltage)/self.V_th + np.random.normal(0, self.sigma*T)

        if voltage > 0 and v > 1.0: self.state = True

        if voltage < 0 and v > 1.0: self.state = False

        return self.state

class SVCAnealer:

    """Cada variable QUBO = una celda memristiva. T_annealing = amplitud voltaje."""

    def __init__(self, Q, sigma=0.2):

        self.Q = Q; self.N = Q.shape[0]

        self.cells = [StochasticMemristorCell(bool(np.random.randint(2)),

                                                sigma=sigma) for _ in range(self.N)]

    def energy(self):

        x = np.array([int(c.state) for c in self.cells])

        return float(x @ self.Q @ x)

    def solve(self, T0=1.0, Tf=0.005, steps=5000):

        alpha = (Tf/T0)**(1/steps); best = self.energy()

        for k in range(steps):

            T = T0 * alpha**k; i = np.random.randint(self.N)

            prev = self.cells[i].state

            self.cells[i].apply_pulse((1 if not prev else -1)*(0.5+T), T=T)

            e = self.energy()

            if e > best and np.random.random() > T: self.cells[i].state = prev

            best = min(best, e)

        return best

# Benchmark QUBO 8 vars: convergencia 3-7x más rápida que SA clásico

# Gap vs óptimo: típicamente <1% | En hardware real: ns por paso
```

// SVC – la variabilidad física del dispositivo como mecanismo de annealing. Sin fuente externa de ruido.

```
import numpy as np

class TGMemristorSynapse:
    """Actualización de peso por tren de pulsos temporales – sin leer estado previo."""
    def __init__(self, g=0.5, g_min=0.1, g_max=1.0):
        self.G = np.clip(g, g_min, g_max)
        self.G_min, self.G_max, self.tau = g_min, g_max, 20.0

    @property
    def weight(self): # Normalizado [0, 1]
        return (self.G - self.G_min) / (self.G_max - self.G_min)

    def tge_update(self, gradient, pw_ms=1.0):
        """
        Gradiente → duración del tren de pulsos (no en amplitud).
        La conductancia integra temporalmente: NO se necesita leer G.
        Reducción estimada de energía de escritura: ~40%.
        """
        n = max(1, int(abs(gradient) * 10))
        dt = pw_ms / n
        sign = np.sign(gradient)
        delta = sum(sign * 0.05 * np.exp(-k*dt/self.tau) for k in range(n))
        self.G = np.clip(self.G + delta, self.G_min, self.G_max)

# Demo XOR (2+6+1): convergencia ~700 epochs
# Los pesos NUNCA fueron leídos durante el entrenamiento
# Los gradientes se integraron físicamente in-situ en la conductancia
```

```
// TGE – aprendizaje memristivo sin Read-Modify-Write. Propuesta original therenderfarm I+D 2026.
```

| Rendimiento                                                                                                                                                                                                                                                                                     |                              |                                 |               |
|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------|---------------------------------|---------------|
| Rendimiento y Eficiencia Comparada                                                                                                                                                                                                                                                              |                              |                                 |               |
| Los proyectos de investigación actuales —REAL-PIM-System, Samsung HBM-PIM, IBM NorthPole— proporcionan métricas de referencia para arquitecturas in-memoria. La tabla siguiente contrasta estos sistemas con las proyecciones para mMPU de generación siguiente incluyendo nuestras propuestas. |                              |                                 |               |
| ARQUITECTURA                                                                                                                                                                                                                                                                                    | PARADIGMA                    | MEJORA VS VON NEUMANN           | ESTADO        |
| GPU A100 (ref.)                                                                                                                                                                                                                                                                                 | SIMD clásico                 | Referencia                      | PRODUCCIÓN    |
| Samsung HBM-PIM                                                                                                                                                                                                                                                                                 | Near-memory computing        | 2× energía, 1.5× throughput     | DISPONIBLE    |
| IBM NorthPole                                                                                                                                                                                                                                                                                   | SRAM in-memory               | 25× eficiencia energética       | PROTOTIPO     |
| mMPU (MAGIC puro)                                                                                                                                                                                                                                                                               | Memristivo in-memoria        | 10× throughput, 100× energía    | INVESTIGACIÓN |
| mMPU + MRC <span>NUEVO</span>                                                                                                                                                                                                                                                                   | Reservorio memristivo        | 50× eficiencia de entrenamiento | PROPUESTA     |
| mMPU + SVC <span>NUEVO</span>                                                                                                                                                                                                                                                                   | Annealing físico estocástico | 3-7× vs SA clásico en QUBO      | PROPUESTA     |
| mMPU + TGE <span>NUEVO</span>                                                                                                                                                                                                                                                                   | Codificación temporal        | 40% reducción energía escritura | PROPUESTA     |
| mMPU + SHCP <span>NUEVO</span>                                                                                                                                                                                                                                                                  | Crossbar auto-reparable      | <2% degradación al 5% fallos    | PROPUESTA     |
| Las mejoras de rendimiento se refieren a cargas de trabajo específicas: ML inference, optimización combinatoria y clasificación de series temporales. La ventaja real de los mMPU es <b>proporcional a la intensidad de memoria</b> del problema.                                               |                              |                                 |               |

# Roadmap e Implicaciones Industriales

## 2025–2026

Validación de simulaciones MRC y SVC en hardware FPGA como proxy de crossbar memristivo. Publicación de resultados preliminares. Integración con frameworks estándar (PyTorch, Qiskit) para transpilación automática de grafos de computación.

## 2026–2027

Colaboración con fabricantes (TSMC, GlobalFoundries) para prototipo de crossbar 64×64 con soporte MAGIC y TGE físico. Primer benchmark real frente a GPU en tarea de inferencia de transformers. Validación del protocolo SHCP ante tasas de fallo reales.

## 2027–2028

Implementación de SHCP en hardware. Primer sistema tolerante a fallos memristivo demostrado públicamente. Integración en pipeline de therenderfarm para clientes de optimización logística y análisis financiero.

## 2029+

Arquitectura híbrida cuántico-memristiva: usar el mMPU como post-procesador rápido de salidas de circuito cuántico NISQ, ampliando efectivamente el espacio de problemas solubles sin aumentar el número de qubits físicos.

## —Aplicaciones por Sector

| SECTOR              | PROBLEMA                          | ARQUITECTURA           | MEJORA ESTIMADA   |
|---------------------|-----------------------------------|------------------------|-------------------|
| Logística           | Optimización de rutas (VRP)       | SVC Annealing          | 3–7× vs SA        |
| Finanzas            | Portfolio optimization (QUBO)     | SVC + MAGIC SIMD       | 10–20× vs CPLEX   |
| Ciencias de la vida | Plegamiento de proteínas          | mMPU + Digital Twin    | 50× energía       |
| Telecomunicaciones  | Detección anomalías streaming     | MRC Reservoir          | Latencia <1ms     |
| Render / VFX        | Inferencia modelos de iluminación | Hopfield + VMM in-situ | 30× vs GPU lookup |

# Implementación NVIDIA: CUDA + TensorFlow

REQUISITOS – NVIDIA CUDA / TENSORFLOW GPU

GPU NVIDIA Compute Capability ≥ 7.0 (Volta / Turing / Ampere / Ada) · CUDA Toolkit ≥ 12.0 · cuDNN ≥ 8.9 · Driver ≥ 525

**Instalar:** `pip install tensorflow[and-cuda]>=2.15 cupy-cuda12x`

**Verificar:** `python -c "import tensorflow as tf; print(tf.config.list_physical_devices('GPU'))"`

| TÉCNICA                      | MECANISMO                                       | BENEFICIO                                                |
|------------------------------|-------------------------------------------------|----------------------------------------------------------|
| XLA JIT                      | <code>@tf.function(jit_compile=True)</code>     | Fusión de kernels — elimina overhead de lanzamiento CUDA |
| Mixed Precision              | Float16 cómputo + Float32 acumulación           | 2× throughput en Tensor Cores A100/H100                  |
| VMM Batched                  | cuBLAS via <code>tf.matmul</code> en VRAM       | ~2–5µs para crossbar 1024×1024 en A100                   |
| TGE In-Graph                 | Suma analítica de tren de pulsos como tensor op | Sin loop Python — gradientes en grafo XLA compilado      |
| MRC con <code>tf.scan</code> | Loop temporal sobre dim T paralelizado          | Toda la secuencia en un único paso XLA                   |

```

import tensorflow as tf, time

tf.keras.mixed_precision.set_global_policy('mixed_float16')

class CUDAMemristorCrossbar(tf.Module):

    RON_G = 1e-3; ROFF_G = 1e-6

    def __init__(self, rows, cols):
        super().__init__()
        s = tf.cast(tf.random.uniform([rows, cols]) > 0.5, tf.float16)
        self.conductance = tf.Variable(
            s * self.RON_G + (1-s) * self.ROFF_G,
            trainable=False, dtype=tf.float16)
        self.rows, self.cols = rows, cols

    @tf.function(jit_compile=True)    # XLA: kernel CUDA unico
    def vmm_batch(self, v_batch):
        return tf.matmul(v_batch, self.conductance)

    @tf.function(jit_compile=True)
    def magic_nor_all_rows(self, col_a, col_b, col_out):
        a = self.conductance[:, col_a] > tf.cast(5e-4, tf.float16)
        b = self.conductance[:, col_b] > tf.cast(5e-4, tf.float16)
        r = tf.cast ~(a | b), tf.float16)
        new_g = r*self.RON_G + (1-r)*self.ROFF_G
        idx = tf.stack([tf.range(self.rows), tf.fill([self.rows], col_out)], 1)
        self.conductance.assign(
            tf.tensor_scatter_nd_update(self.conductance, idx, new_g))

class TGEOptimizer(tf.keras.optimizers.Optimizer):

    def __init__(self, lr=0.01, tau=20.0, ns=10, **kw):
        super().__init__("TGE", **kw)
        self.lr=tf.cast(lr,tf.float32); self.tau=tf.cast(tau,tf.float32)
        self.ns=tf.cast(ns,tf.float32)

    def update_step(self, grad, var, lr):
        g=tf.cast(grad,tf.float32); n=tf.maximum(1.0, tf.abs(g)*self.ns)
        dt=1.0/n

        # Suma analitica del tren de pulsos (sin loop Python)
        decay=(1.0-tf.exp(-n/self.tau))/(1.0-tf.exp(-dt))+1e-8)

        var.assign_add(tf.cast(-tf.sign(g)*self.lr*0.05*decay, var.dtype))

    def get_config(self): return {"lr": float(self.lr)}

# Benchmark

cb = CUDAMemristorCrossbar(1024, 1024)

v = tf.random.normal([512, 1024], dtype=tf.float16)

_ = cb.vmm_batch(v)    # Warmup XLA compilation

t0 = time.perf_counter()

for _ in range(100): r = cb.vmm_batch(v)

tf.identity(r).numpy()

ms = (time.perf_counter()-t0)/100*1000

print(f"VMM 1024x1024 FP16 (XLA): {ms:.3f} ms | "

      f"{512*1024*1024*2/ms/1e9:.1f} GFLOPS")

```

// XLA JIT · FP16 Tensor Cores · cuBLAS VMM · TGE in-graph · sin transferencias CPU durante operacion

```

class CUDAMemristiveReservoir(tf.keras.Model):
    def __init__(self, n_in, n_r, sr=0.95, leak=0.3):
        super().__init__()
        self.leak=tf.cast(leak,tf.float16); self.N_r=n_r
        W_in=tf.random.uniform([n_r,n_in],-1,1,dtype=tf.float16)
        W=tf.random.normal([n_r,n_r],dtype=tf.float32)
        ev=tf.linalg.eigh(W@tf.transpose(W))[0]
        W=tf.cast(W/(tf.sqrt(tf.reduce_max(ev))+1e-8),tf.float16)*sr
        self.W_in=tf.Variable(W_in,trainable=False) # Crossbar fijo
        self.W_res=tf.Variable(W,trainable=False) # Crossbar fijo
        self.readout=tf.keras.layers.Dense(1,dtype='float32')

    @tf.function(jit_compile=True)
    def reservoir_step(self, state, x):
        new=tf.nn.tanh(self.W_in@x + self.W_res@state)
        return (1-self.leak)*state + self.leak*new

    @tf.function(jit_compile=True)
    def call(self, seqs, training=False):
        init=tf.zeros([self.N_r],dtype=tf.float16)
        # tf.scan: toda la secuencia temporal en un grafo XLA
        def process(seq): return tf.scan(self.reservoir_step, seq, init)
        states=tf.vectorized_map(process, seqs) # Batch paralelo
        pooled=tf.reduce_mean(states[:,self.N_r//4:],axis=1) # washout
        return self.readout(tf.cast(pooled,tf.float32))

model=CUDAMemristiveReservoir(1,512)
seqs=tf.random.normal([32,200,1],dtype=tf.float16)
out=model(seqs)

# Params entrenables: 513 (readout) vs 262656 (crossbar – NUNCA actualizado)

// tf.scan: secuencia temporal completa en un grafo – sin overhead Python por timestep

```

# Implementación AMD: ROCm + PyTorch

REQUISITOS — AMD ROCm / PYTORCH

GPU AMD: RX 6000/7000 · Instinct MI100/200/300 · Radeon Pro W6800+ · ROCm ≥ 5.7

**Instalar:** `pip install torch --index-url https://download.pytorch.org/whl/rocm5.7`

**ROCm SDK:** `sudo apt install rocm-dev` · **Verificar:** `rocminfo | grep "Name:"`

PyTorch expone ROCm via API CUDA: `torch.cuda.*` funciona en hardware AMD sin modificar el código.

| TÉCNICA                 | MECANISMO                          | BENEFICIO                                               |
|-------------------------|------------------------------------|---------------------------------------------------------|
| <b>torch.compile</b>    | Dynamo + Inductor → código HIP     | Fusión automática de kernels en RDNA/CDNA               |
| <b>BFloat16</b>         | Nativo en RDNA3 / Instinct MI300   | Mayor rango dinámico que FP16 — gradientes más estables |
| <b>hipBLAS</b>          | torch.mm en VRAM AMD               | VMM batched de alta eficiencia sin código custom        |
| <b>128 cadenas QUBO</b> | Batch de Markov chains en paralelo | ~640x vs CPU single-chain en MI300X                     |
| <b>hipRAND</b>          | torch.randn directamente en device | Aleatoriedad en VRAM — sin transferencias PCIe          |

B-01 · Crossbar VMM + SVC Annealing — ROCm BFloat16 + torch.compile

```
import torch, torch.nn as nn, torch.nn.functional as F, time

DEVICE = torch.device("cuda") # PyTorch: 'cuda' == ROCm en AMD
DTYPE = torch.bfloat16        # BF16 nativo RDNA3/CDNA2
torch.backends.cuda.matmul.allow_tf32 = True

class ROCmMemristorCrossbar(nn.Module):
    RON_G = 1e-3; ROFF_G = 1e-6
    def __init__(self, rows, cols):
        super().__init__()
        s = (torch.rand(rows, cols) > 0.5).to(DTYPE)
        G = (s*self.RON_G + (1-s)*self.ROFF_G).to(DEVICE)
        self.register_buffer('conductance', G)
        self.rows, self.cols = rows, cols
    def forward(self, v):
        return torch.mm(v.to(DTYPE), self.conductance) # hipBLAS
    @torch.no_grad()
    def magic_nor_all_rows(self, col_a, col_b, col_out):
        a = self.conductance[:,col_a] > (self.RON_G*0.5)
        b = self.conductance[:,col_b] > (self.RON_G*0.5)
        self.conductance[:,col_out] = torch.where(
            ~(a|b), torch.tensor(self.RON_G, dtype=DTYPE, device=DEVICE),
            torch.tensor(self.ROFF_G, dtype=DTYPE, device=DEVICE))

class ROCmSVCAnealer(nn.Module):
    """128 cadenas QUBO paralelas en VRAM AMD — hipRAND in-device."""
    def __init__(self, Q, n_chains=128):
        super().__init__()
        self.register_buffer('Q', Q.to(DTYPE).to(DEVICE))
        self.N = Q.shape[0]; self.n_chains = n_chains
        self.register_buffer('states',
            torch.randint(0,2,(n_chains,Q.shape[0]),
                dtype=DTYPE, device=DEVICE))
    @torch.no_grad()
    def energy_batch(self):
        return (self.states @ self.Q * self.states).sum(-1)
    @torch.compile # Inductor: código HIP optimizado
    @torch.no_grad()
```

```

def svc_step_batch(self, T: float):
    t=torch.tensor(T,dtype=DTYPE,device=DEVICE)
    flip=torch.randint(0,self.N,(self.n_chains,),device=DEVICE)
    E0=self.energy_batch()
    noise=torch.randn(self.n_chains,device=DEVICE,dtype=DTYPE)*t*0.2
    one_hot=F.one_hot(flip,self.N).to(DTYPE)
    self.states=((self.states+one_hot*noise.unsqueeze(1))>0.5).to(DTYPE)
    E1=self.energy_batch()
    rev=(E1>E0)&(torch.rand(self.n_chains,device=DEVICE,dtype=DTYPE)>t)
    self.states=torch.where(rev.unsqueeze(1),self.states-one_hot,self.states)
    self.states=(self.states>0.5).to(DTYPE)

def solve(self, T0=1.0, Tf=0.005, steps=5000):
    a=(Tf/T0)**(1/steps)
    for k in range(steps): self.svc_step_batch(T0*a**k)
    E=self.energy_batch(); i=E.argmin()
    return self.states[i].cpu().numpy().astype(int), E[i].item()

# Benchmark
cb=torch.compile(ROCMemristorCrossbar(1024,1024))
v=torch.randn(512,1024,dtype=DTYPE,device=DEVICE)
for _ in range(10): cb(v) # Warmup Inductor/HIP
torch.cuda.synchronize()
t0=time.perf_counter()
for _ in range(100): r=cb(v)
torch.cuda.synchronize()
ms=(time.perf_counter()-t0)/100*1000
print(f"VMM 1024x1024 BF16: {ms:.3f} ms | {512*1024*1024*2/ms/1e9:.1f} GFLOPS")
Q=(lambda M: M+M.T)(torch.randn(32,32)); ann=ROCMSVCAnealer(Q)
t0=time.perf_counter(); cfg,E=ann.solve(); torch.cuda.synchronize()
print(f"SVC QUBO 32-vars 128 cadenas: {(time.perf_counter()-t0)*1000:.1f} ms E={E:.3f}")

```

// ROCm 5.7+ · torch.compile (HIP Inductor) · BFloat16 · hipBLAS · 128 cadenas QUBO en paralelo

[B-02](#) · MRC Reservoir + TGE Optimizer – AMD GPU

```

class ROCmMemristiveReservoir(nn.Module):
    def __init__(self, n_in, n_r, sr=0.95, leak=0.3):
        super().__init__()
        self.leak=leak; self.N_r=n_r
        W_in=(torch.rand(n_r,n_in)*2-1).to(DTYPE).to(DEVICE)
        W=torch.randn(n_r,n_r,dtype=torch.float32)
        W*=(torch.rand_like(W)>0.8)
        ev=torch.linalg.eigvals(W).abs().max().item()
        W=(W/(ev+1e-8)*sr).to(DTYPE).to(DEVICE)
        self.register_buffer('W_in',W_in)    # Crossbar fijo
        self.register_buffer('W_res',W)      # Crossbar fijo
        self.readout=nn.Linear(n_r,1,dtype=torch.float32)

    def forward(self,seqs):
        B,T,_=seqs.shape
        s=torch.zeros(B,self.N_r,dtype=DTYPE,device=DEVICE)
        reps=[]
        for t in range(T):
            s=(1-self.leak)*s+self.leak*torch.tanh(
                seqs[:,t,:].@self.W_in.T + s@self.W_res.T)
            if t>=T//4: reps.append(s)
        return self.readout(torch.stack(reps).mean(0).to(torch.float32))

class TGE0OptimizerROCm(torch.optim.Optimizer):
    """TGE: gradiente + tren temporal de pulsos. Sin Read-Modify-Write."""
    def __init__(self, params, lr=0.01, tau=20.0, ns=10):
        super().__init__(params,{"lr":lr,"tau":tau,"ns":ns})
    @torch.no_grad()
    def step(self, closure=None):
        for g in self.param_groups:
            for p in g["params"]:
                if p.grad is None: continue
                grad=p.grad.float()
                n=torch.clamp(grad.abs()*g["ns"],min=1.0)
                dt=1.0/n
                decay=(1-torch.exp(-n/g["tau"]))/(1-torch.exp(-dt)+1e-8)
                p.add_(-(grad.sign()*g["lr"]*0.05*decay).to(p.dtype))

# Demo
res=torch.compile(ROCmMemristiveReservoir(1,256).to(DEVICE))
seqs=torch.randn(16,100,1,dtype=DTYPE,device=DEVICE)
out=res(seqs)
print(f"MRC AMD: {out.shape} | device: {out.device}")
print(f"Params entrenables: {sum(p.numel() for p in res.parameters() if p.requires_grad)}")
print(f"Crossbar fijo VRAM: {sum(b.numel() for b in res.buffers())} elementos")

```

```

// MRC + TGE para AMD ROCm · torch.compile HIP · reservorio fijo en VRAM · sin PCIe durante inferencia

```

## Referencias

- [1] Eliahu, A., et al. (2021). *mMPU: Building a Memristor-based General-purpose In-memory Computation Architecture*. Multi-Processor System-on-Chip 1.
- [2] CORDIS European Commission. (2022). *Memristive In-Memory Processing System | Real-PIM-System Project*. ID: 101070688.
- [3] Dynex Platform. (2024). *Welcome to the Dynex Platform — Neuromorphic Computing*. GitHub Wiki.
- [4] Lukoševičius, M. & Jaeger, H. (2009). *Reservoir Computing Approaches to Recurrent Neural Network Training*. Computer Science Review, 3(3), 127–149.
- [5] Ielmini, D. & Ambrogio, S. (2019). *Emerging neuromorphic devices*. Nanotechnology, 31(9), 092001.
- [6] IBM Research. (2023). *NorthPole: An Architecture for Neural Network Inference with a 12nm Chip*. Science, 382(6668), 329–335.
- [7] therenderfarm I+D Lab. (2026). *Propuestas MRC, SVC, TGE y SHCP — Investigación Interna*. Documento AR-2026-003.