



## PARTE I

Neste documento, parte I, serão referidos os princípios físicos que levaram ao armazenamento permanente de informação digital em dispositivos microeletrônicos de estado sólido, designados por memórias não voláteis, que designámos por “Gaiola Eletrostática de Bits”, especificamente as que são baseadas no transistor MOSFET de porta flutuante condutora ou isoladora.

A empresa pioneira nesta tecnologia, INTEL, foi formada para desenvolver e explorar o mercado destes circuitos que iriam revolucionar os sistemas computacionais e instrumentos científicos eletrónicos que hoje conhecemos. Até 1980, o mercado das memórias foi a principal fonte de rendimento da INTEL.

Os métodos de programação destes dispositivos são os de injeção de portadores quentes, HCI, até 1986, e depois, de tunelamento quântico microscópico e HCI, que foram fundamentais para o desenvolvimento das arquiteturas NAND e NOR das memórias *Flash*.

O documento salienta as inovações realizadas nos primeiros 30 anos desta tecnologia, que criou muitas empresas e estabeleceu um negócio do mercado global que, em 2025, valia cerca de 70 mil milhões de euros, quase equivalente ao dos processadores (100 mil milhões).

Com a redução da dimensão dos transístores para o mínimo atualmente possível, 20 nm x 20 nm<sup>1</sup>, o aparecimento de novas tecnologias de integração monolítica tridimensional das memórias V-NAND, permitiu que a célebre Regra de Moore <sup>2</sup>permanecesse válida.

Neste documento, parte I, são referidas as fragilidades deste tipo de memórias e as várias causas que podem levar à perda dos dados armazenados, mas, na Parte II, são referidas as técnicas que permitem que o controlador associado permita tornar estas memórias fiáveis com estratégias inteligentes do seu uso, técnicas de correção de erros, de modo a tornarem-se bastante fiáveis quer no formato *pen* USB quer no formato *Solid State Disc*, SSD, nas versões *consumere enterprise*. Estas técnicas permitiram fazer memórias num *chip* com 1TeraByte em 2019 e 8 TeraByte em 2025

Em 2025, o negócio das memórias não voláteis teve uma elevada rentabilidade, variável entre 55% a 85%, sendo liderado pelas empresas Samsung Electronics, SK Hynix, Micron e Kioxia.

Em 2026, com o crescimento dos sistemas de Inteligência Artificial, IA, e dos *Data Center*, que requerem enormes capacidades de memória RAM extremamente rápida, os fabricantes estão a mudar capacidade das fábricas para a produção das memórias RAM *High bandwidth Memory*, HBM, pelo que os custos das memórias RAM convencional e *Flash* devem subir muito devido a escassez de produção.

### Os princípios das memórias de estado sólido

Em 1897, Robert [Wood](#) (1868 - 1955), descobriu que uma ponta metálica muito próxima de um cátodo, no vácuo, à temperatura ambiente, se fosse submetida a uma tensão muito elevada provocava uma corrente elétrica no vácuo, que era considerado um excelente isolador elétrico. Este fenómeno é, hoje em dia, considerado como “emissão fria de eletrões”.

---

<sup>1</sup> Este mínimo só é possível com a máquina mais complexa feita pelos humanos, a [máquina ASML de litografia EUV com luz de 13 nm](#) de comprimento de onda.

<sup>2</sup> - Com o mesmo custo, o nº de transístores por circuito com a mesma área duplicaria em cada 2 anos.



Em 1925, [Robert Millikan](#) (1868 - 1953), o cientista da experiência da carga do elétron que lhe deu o prêmio Nobel em 1923, e [Carl Eyring](#) (1889 - 1951) revisitaram a experiência de emissão fria de Robert [Wood](#). Millikan e Eyring fizeram gráficos experimentais que relacionavam linearmente a corrente com o inverso da tensão aplicada.

Os elétrons saíam do filamento, ou da ponta metálica, com uma energia muito inferior à exigida pelo trabalho de extração eletrônica do condutor, mas não conseguiram explicar o fenômeno.

Mais tarde, em 1928, Millikan e [Charles Lauritsen](#) (1862 – 1968) estudaram o relacionamento entre a emissão termiônica e a emissão de elétrons frios e concluíram que não existe interação, sendo dois fenômenos completamente independentes. Publicaram o artigo “[Relations of Field Currents to Thermionic Currents](#)” na *Physical Review*, vol 14 n.8, onde reuniram todos

os dados que a equipa tinha obtido experimentalmente desde 1925 e desenvolveram expressões empíricas que são válidas para quando os dois efeitos estão simultaneamente presentes. Estes investigadores não conseguiram explicar este comportamento, mas o gráfico  $I(V)$ , Fig. 1, que obtiveram, ficou conhecido por gráfico de Millikan-Lauritsen.

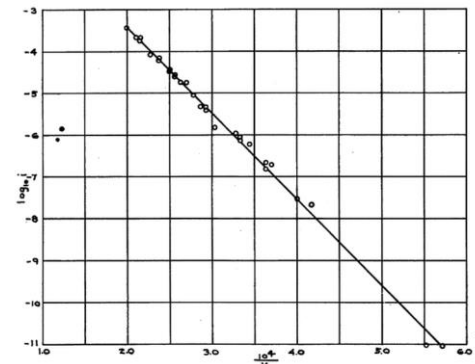


Fig. 1- O célebre gráfico de Millikan Lauritsen

A explicação científica da “emissão fria” baseia-se na teoria da mecânica quântica e o fenômeno chama-se [tunelamento quântico microscópico](#), que explica os mecanismos de condução através de isoladores elétricos finos. Foi em 1927, que [Ralph. H. Fowler](#) (1889 - 1944) e [Lothar Nordheim](#) (1899 - 1985) explicaram, usando as teorias da mecânica quântica, que campos elétricos elevados aplicados a um isolador em percursos muito pequenos, podem originar transporte de cargas elétricas ou seja, [a condução elétrica através dos isoladores](#).

Nos anos 70, houve muita investigação sobre as teorias da injeção de portadores minoritários nos semicondutores nomeadamente quando estes adquirem muita energia (equivalente a terem uma temperatura muito elevada). O fenômeno designa-se por *Hot Carrier Injection* e tem um efeito de desgaste sobre os semicondutores, contribuindo para a sua degradação. Também o tunelamento eletrónico têm um efeito de desgaste das propriedades dos isoladores dos semicondutores. Ambos os fenômenos são a base do funcionamento das memórias Flash, usando o conceito de transistor MOSFET de porta flutuante. Estas memórias têm vários defeitos devidos ao desgaste dos transístores, mas têm sido desenvolvidas muitas técnicas tendo por objetivo minorar estes defeitos e tornar as memórias relativamente fiáveis.

### O Transístor de Porta ou “Gate” flutuante

O conceito teórico de usar carga armazenada para controlar transístores para memorizar informação foi patenteado, separadamente, pela primeira vez, por [Jack Morton](#) ( ...- 1971) e [Ian Ross](#) (– 2013), , nos Laboratórios Bell<sup>3</sup>, mas ainda não havia os transístores adequados para isto – o MOSFET.

Em 1959, o transistor MOSFET foi inventado pelos engenheiros [Dawon Kahng](#) (1931 - 1992) e [Mohamed Atalla](#) (1924 - 2009) nos Bell Labs, tendo o primeiro protótipo funcional sido apresentado ao público em 1960

Em 1961, [Chih-tang Sah](#) (1932 – 2025) na Fairchild lançou a ideia de que um transistor MOSFET poderia ser usado para reter a carga no condutor da porta de comando durante vários dias.

<sup>3</sup> Ian Ross foi o proponente do método muito importante, na eletrónica, do crescimento epitaxial sobre silício.



Em 1966, [John Szedon](#), [Edgar A. Sack](#) (1930 - ...) , [Ting L. Chu](#) (1924 -2017) e outros, na Westinghouse Microelectronics, desenvolveram uma estrutura teórica do transistor MOSFET MNOS (com camadas Metal-Nitreto-Óxido-Silício) em que a porta flutuante seria feita com um isolador, o Nitreto de Silício,  $\text{Si}_3\text{N}_4$ , que ficou designado por Transistor de armadilha “*Charge Trap Transistor*”, CTT.

A porta flutuante de um transistor, quando realizada com nitreto de silício, tem a capacidade de reter a carga elétrica de forma localizada em pequenos pontos (armadilhas), contrariamente ao que acontece com uma porta condutora de metal ou de polissilício, em que a carga se distribuiria uniformemente pela superfície do condutor.

Em 1967, [Dawon Kahng](#) (1931 - 1992) e [Simon Sze](#) (1936 - 2023), dos Laboratórios Bell, desenvolveram a teoria de como a porta flutuante de um dispositivo semiconductor MOSFET poderia ser usado como célula de uma ROM (Read Only Memory) reprogramável. Daqui em diante apelidaremos este transistor por Transistor de Porta Flutuante, TPF.

Ainda em 1967, [H.A. Richard Wegener](#), na empresa [Sperry](#) transformou o conceito teórico do transistor de armadilha de carga e realizou o primeiro transistor funcional. Esta estrutura viria a ser melhorada em 1977, pelo transistor [SONOS](#) por PCY Chen da [Fairchild Camera and Instruments](#).

É um transistor MOSFET canal N, num substrato P, seguido de várias camadas: 2 nm de óxido fino, 5 nm de nitreto de silício, 5 a 10 nm de óxido de silício e a porta de controlo em Polissilício.

Esta estrutura de transistor pode funcionar como memória não volátil (nas memórias Flash do tipo *Charge Trap*), substituindo a porta flutuante de polissilício de um transistor. Contudo, a tecnologia foi considerada, na época, inferior à do transistor de *Floating Gate* condutora, e permaneceu dedicada para usos em nichos de aplicações, devido às altas tensões exigidas (perto de 50V). Todavia, esta tecnologia viria a ser fundamental nas memórias Flash realizadas a partir de 2002.

Na Fig. 2, a porta flutuante, PF, do transistor MOSFET de canal N, está eletricamente isolada, e está separada dos outros condutores por óxido de silício que é um excelente isolador elétrico. A porta flutuante é a gaiola ou onde serão guardados os eletrões. No transistor de armadilha de carga, a porta do transistor é feita de material isolante e é onde serão guardadas as cargas elétricas.

### O modo de gravação do Transistor TPF condutora

Por [tunelamento quântico microscópico ou por injeção de portadores quentes](#), uma tensão positiva, muito elevada, aplicada à porta P de controlo, durante um certo tempo, pode injetar uma carga elétrica, constituída por um certo número de eletrões, na porta flutuante do TPF. Nesta fase, designada por fase de gravação, ou de programação, os eletrões virão do substrato, e atravessam a camada de óxido fino do transistor.

Como as fugas são muito pequenas, esta carga permanecerá armazenada na porta flutuante do transistor.

Controlando o tempo de duração da tensão elevada aplicada à porta P, de controlo, pode controlar-se o nível de carga armazenada na porta flutuante, Fig. 3a).

Invertendo a polaridade da tensão entre a porta e o substrato os eletrões armazenados na PF podem regressar ao substrato ficando o transistor sem carga armazenada na porta flutuante. Esta fase designa-se por apagamento da célula de memória.

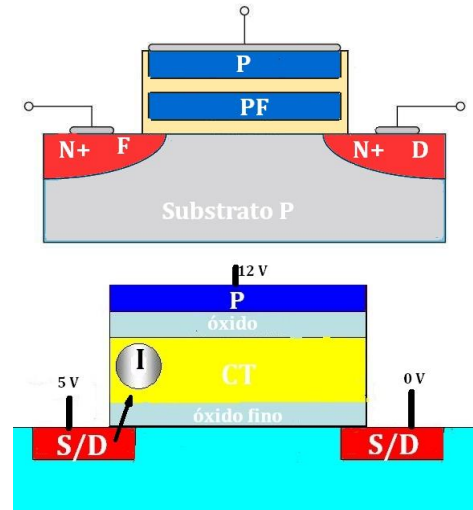


Fig. 2 -Transistores MOSFET de PF e de CT.



Quer o mecanismo de programação quer o de apagamento da carga pode ser feito por tunelamento quântico microscópico ou por injeção de portadores quentes HCI (*Hot Carrier Injection*). Ambos os mecanismos introduzem alguma degradação do transistor sendo que a HCI é mais rápida, [mas desgasta mais o transistor](#).

### O controlo da injeção de carga na PF do transistor

Na Fig. 4 representam-se 4 níveis diferentes de carga armazenada na porta flutuante, provocados pelo campo elétrico  $E_p$ , suficiente para produzir tunelamento quântico, ou por Injeção de portadores quentes, HCI, e carregar a porta flutuante do transistor com 4 níveis diferentes da carga  $Q$ , através da duração do pulso  $T_p$ , de programação do TPF.

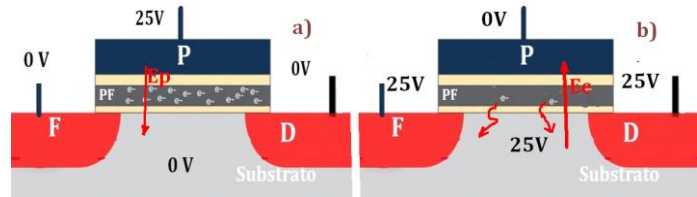


Fig. 3 - Como carregar e descarregar a Porta flutuante do TPF.

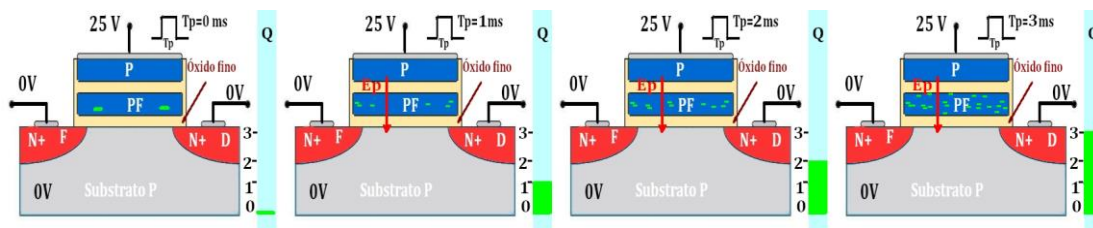


Fig. 4 - Carga armazenada na porta flutuante de um TPF em função da duração temporal do pulso de programação.

Pode facilmente pensar-se que a carga analógica  $Q$ , com 4 níveis analógicos distintos, poderia representar a informação digital de 2 bit armazenados num único transistor.

Atualmente, são produzidas memórias digitais baseadas no TPF que apenas guardam um bit por transistor (célula), designadas por **SLC** “*Single Level Cell*”, ou que guardam 2 bit por transistor designadas por **MLC** “*Multi Level Cell*” ou, ainda, por transistores que guardam 8 níveis de carga correspondentes a 3 bit digitais, designadas por **TLC** “*Triple Level Cell*”.

Hoje em dia, as células das memórias de TPF mais usuais são as TLC usadas para fabricar os “discos” que são designadas por SSD (*Solid State Disk*) e que visam ser alternativa aos discos magnéticos rígidos *Hard Disk*, HD, dos computadores. Para memórias baratas tipo *Pen* (flash) são usadas células **QLC** com 16 níveis de carga, ou seja, representam 4 bits digitais. A evolução tecnológica permite esperar que, para o ano de 2030, estejam disponíveis as memórias **PLC** (*Penta Level Cell*) em que serão guardados 32 níveis analógicos de carga elétrica, correspondentes a 5 bit digitais.

### A leitura da carga armazenada no TPF

No modo de leitura do TPF, a tensão aplicada à porta de controlo,  $P$ , é muito pequena quando é comparada com a tensão usada no modo de gravação, não havendo qualquer injeção significativa de cargas por tunelamento quântico microscópico ou por injeção de portadores quentes.

Quando a tensão da porta de controlo,  $P$ , é nula, a carga elétrica armazenada na porta flutuante cria um campo elétrico,  $E_q$ , permanente, entre o substrato do transistor e a porta flutuante. Este campo elétrico determina a tensão mínima de controlo, limiar de condução do transistor, acima designada por tensão de *Threshold*,  $V_t$ .

A tensão  $V_t$  depende do nível de carga armazenada na porta flutuante condutora e vai influenciar a característica de condução do transistor, Fig. 5.

Normalmente os transistores de porta flutuante, sem carga elétrica armazenada, têm uma tensão limiar de controlo que é muito próxima de 0 V.



Esta tensão designa-se por *Threshold Voltage*,  $V_t$ . Se a tensão entre Porta e Fonte,  $V_{pf}$ , for  $V_{pf} > V_t$ , o transistor conduz e estará ao corte (não conduz) quando for  $V_{pf} < V_t$ .

A aplicação de uma tensão positiva na porta de controlo vai criar um campo elétrico  $E_p$  de sinal contrário ao campo resultante da carga armazenada do transistor modificando o valor da tensão  $V_t$ .

Na Fig. 6 podem observar-se 4 transístores com diferentes níveis de carga acumulada submetidos a uma tensão de controlo,  $V_P$ , de 2,5 V.

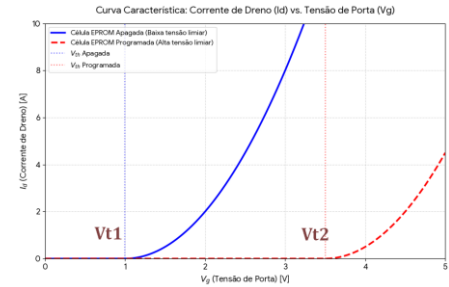


Fig.5 - Alteração das características do TPF com a carga armazenada.

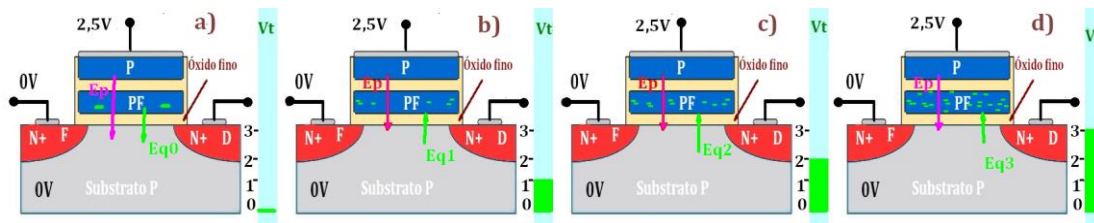


Fig. 6 – A carga armazenada na porta flutuante de um TPF determina o valor da tensão limiar de controlo,  $V_t$ , do TPF.

Na Fig.6, os transístores a), b) e c) estariam no estado de condução porque a tensão de 2,5 V é superior às suas tensões de limiar  $V_t$ . O transistor d) estaria cortado pois  $2,5 \text{ V} < 3 \text{ V}$ .

Assim, o processo de leitura da carga existente na porta flutuante do transistor consiste em aplicar uma rampa de tensão, que cresce linearmente no tempo, a partir de 0 V. A evolução desta rampa de tensão vai obrigar o estado do transistor a mudar da situação de corte para a de condução.

O estado de condução do transistor é determinado por um sensor de corrente associado ao terminal Dreno, a que está aplicada uma tensão positiva. Na Fig. 6, quando a rampa tem um valor menor do que  $V_t = 1 \text{ V}$ , Fig 6a) o transistor estará no estado OFF e a carga elétrica armazenada é considerada nula. Assim que a rampa ultrapassa o valor de 1 V o transistor terá uma carga  $Q_1$  e passará para o modo de condução. Assim que  $V_P > 2 \text{ V}$  a carga  $Q_2$  será determinada e os transístores a) b) e c) estão ON e d) estará OFF. É este o processo correntemente usado para avaliar o valor da carga armazenada em cada transistor de porta flutuante.

### Escrita e leitura do transistor de armadilha de carga

Através do controlo adequado das tensões de porta, dreno e fonte pode armazenar-se ou retirar cargas da porta flutuante isoladora, do transistor de armadilha de carga.

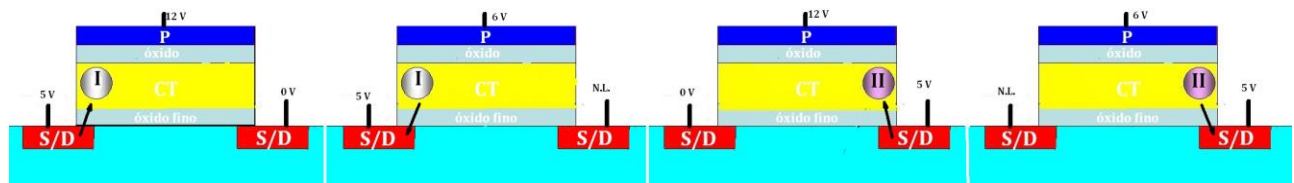


Fig. 7 - Gravação e apagamento de cargas no Transistor de Armadilha de carga.

As cargas elétricas ficam posicionadas na porta isolante junto dos elétrodos (S ou D) mais positivos, podendo a porta ter cargas diferentes nos dois pontos, Fig. 7, uma vez que elas não se misturam por estarem em regiões diferentes do isolador. O transistor até pode armazenar mais do que um nível de carga, como se mostra na Fig. 8 com a armadilha I e a armadilha II, localizadas em diferentes pontos.





Com os transístores de armadilha de carga pode gravar-se mais do que um bit em cada transístor, porque a porta flutuante não é condutora e as cargas ficam alojadas em sítios bem determinados sem se misturarem.

Na Fig.8 pode ver-se uma célula de cargas armadilhadas que tem duas regiões de carga I e II, que podem ser lidas individualmente por uma sequência de tensões aplicadas à Porta de controlo P e à Fonte e ao Dreno. Foram, até, já desenvolvidos métodos para armazenar quatro níveis de cargas em quatro regiões distintas da porta isoladora.

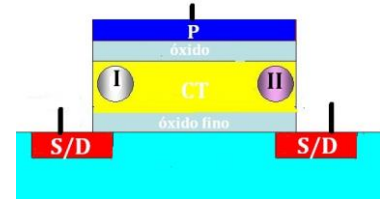


Fig. 8 - Célula CTT, com 2 cargas distintas armazenadas na mesma porta isoladora.

### O uso do óxido de Háfnio em vez de óxido de Silício

Na década de 90 os TPF tinham dimensões grandes e acumulavam cargas que podiam conter 10 000 a 100 000 eletrões.

Atualmente, em 2026, com a miniaturização dos transístores ao nível de 20 nm x 20 nm, de área, a capacidade de armazenamento de carga é muito pequena. Nestes transístores, se fosse usado óxido de silício como isolador, com a espessura de 2 nm, a capacidade da porta do TFG seria extremamente pequena. Teria o valor de 6,2 aF (attoFarad –  $1\text{aF} = 10^{-18}\text{ F}$ ). Se a porta do transístor fosse carregada com apenas 3,3 V, durante a fase de programação, a carga Q armazenada ( $Q = C_p \times V$ ) seria de cerca de 22,8 aC (attoCoulomb) que corresponde apenas a ter 143 eletrões armazenados na porta flutuante do transístor.

Na indústria, para as células mais usuais do tipo TLC (Triple Level Cell), considera-se como mínimo aceitável a existência de 500 eletrões armazenados na porta, pelo que em vez de óxido de silício é usado o óxido de Háfnio<sup>4</sup> que é quase tão bom isolador como o óxido de silício mas tem uma constante dielétrica, K, de cerca de 20 a 25, que é 5 a 6 vezes maior que a do óxido de silício ( $k = 3,9$ ). Usando o óxido de háfnio para a mesma espessura de óxido de 2 nm o transístor já armazenaria cerca de 700 a 850 eletrões. Alternativamente, os 500 eletrões armazenados na porta flutuante poderiam ser obtidos com uma espessura maior do óxido fino.

A tensão de programação para os transístores de porta flutuante, atualmente mais usada, é de 12,5 V. Com esta tensão, o tempo necessário para injetar 500 eletrões na PF é bastante pequeno, cotribuindo ara a diminuição do tempo de programação.

Em 2007, a INTEL fez a passagem do processo de fabricação de transístores de 65 nm, para o de 45 nm, tendo introduzido o óxido de háfnio como isolador. A mesma capacidade da porta dos transístores passou a poder ser obtida com uma maior espessura do óxido fino. Isto tinha como objetivo principal conseguir reduzir os efeitos de tunelamento que originariam correntes parasitas nos transístores dos processadores Penryn e Xéon (que seriam realizados no processo de 45 nm).

A introdução do háfnio é considerada a maior mudança no processo de fabricação de transístores, depois dos anos 60. O Háfnio permitiu que se chegasse aos atuais transístores FinFET e AGA (All Gate Around).

### A resolução da carga dos TPF Multinível

Considerando o valor mínimo da carga máxima armazenada na porta de um TPF de 500 eletrões, na célula TLC (8 níveis, 3 Bit digitais), o sistema de gravação da carga elétrica terá de ajustar os níveis da carga armazenada, para cada transístor, com a separação entre níveis de cerca de 70 eletrões, isto é: 0, 70, 140, 210, 280, 350, 420 e 490 eletrões.

<sup>4</sup> O Háfnio é uma daquelas terras raras tão desejadas na indústria dos semicondutores. Não existe livre na natureza, mas está associado, em pequenas percentagens, ao zircónio, que é mais abundante, mas não é fácil separá-los. Também é usado no controlo de reatores nucleares.



Este processo é feito através do controlo iterativo do tempo de pulso da tensão de programação aplicada à porta de controlo do TPF e verificando o estado de condução do transistor, através da observação da tensão  $V_t$ , limiar de condução, que depende da carga armazenada e da medida da corrente de Dreno.

Este processo é lento, pois é iterativo e exige a repetição da sequência de programação / verificação da carga armazenada, para acertar o nível analógico da carga pretendida, correspondente à informação digital desejada.

### A degradação temporal da carga armazenada no TPF

Os óxidos de silício e de háfnio não são isoladores perfeitos. A PF, “gaiola”, onde são armazenados os eletrões tem algumas perdas, que podem ser minimizadas aumentando a espessura dos óxidos que separam os vários transístores da memória, mantendo, contudo, a espessura do óxido fino, para que o tunelamento ou a injeção de eletrões quentes seja possível.

Obviamente, a temperatura do transistor tem efeito pois sabe-se que um aumento de cerca de 5 a 10° C duplica as correntes de fuga nos semicondutores. Atualmente, à temperatura ambiente de 25 °C, a degradação temporal prevista para uma célula TLC (3 bit), não eletricamente alimentada, é de cerca de 5 anos, ou seja, estima-se uma perda máxima de 35 eletrões durante cinco anos.

O aumento da temperatura da célula em repouso, para 35 °C pode significar a alteração da duração da informação guardada para apenas um ou dois anos.

### A degradação provocada pela luz UV de apagamento da célula

Inicialmente, em 1971, quando foram usados transístores PMOS normais para época, com 30 um x 30 um, mas que seriam gigantes nos dias de hoje, a previsão da duração da informação (carga) guardada era de cerca de 20 a 30 anos. Nesta época, para apagar a gravação registada nos transístores recorria-se a luz ultravioleta com cerca de 200 nm de comprimento de onda; a luz era aplicada durante cerca de 20 minutos, destruía as cargas, mas também degradava o semiconductor pelo que os transístores ao fim de 10 a 20 operações de apagamento deixavam de funcionar.

Em 1975, a INTEL introduziu o processo NMOS (substrato P) e os transístores de porta flutuante passaram a durar cerca de 100 a 1000 sessões de apagamento com luz UV.

### A degradação provocada por sucessivas gravações

Os fenómenos de tunelamento quântico e de injeção de portadores quentes, usados na operação da gravação ou de apagamento de uma célula de memória baseada no TPF, provocam efeitos nefastos cumulativos. Basicamente, com o uso continuado destes processos, são criados eletrões que ficam aprisionados permanentemente no dielétrico e que alteram as características do transistor de modo definitivo.

Sempre que uma célula perde eletrões, o seu limiar de condução,  $V_t$ , desce, podendo entrar na zona do estado digital vizinho, causando um erro de leitura. Em sentido contrário, ler repetitivamente uma célula pode injetar acidentalmente uma quantidade mínima de carga numa célula vizinha, aumentando o seu limiar de condução,  $V_t$ .

Em 2026, o número de ciclos úteis (com degradação aceitável) de gravação / apagamento de células baseadas em TPF que são usualmente “garantidos” para as diferentes células multinível, que demonstram a sua resiliência, tipicamente são os representados na Tab. 1.

| Célula | Níveis de carga | Milhares de ciclos |
|--------|-----------------|--------------------|
| SLC    | 2               | 50 a 100           |
| DLC    | 4               | 3 a 10             |
| TLC    | 8               | 1 a 3              |
| QLC    | 16              | 0,5 a 1            |
| PLC    | 32              | 0,1 a 0,5          |

Tab. 1- Resiliência das células dos TPF.



## O nascimento da empresa INTEL

Na primavera de 1968, [Gordon Moore](#) (1929 – 2023) apareceu na casa de [Bob Noyce](#) (1927 – 1990) e Moore sugeriu que a recente memória semicondutora, uma tecnologia emergente, poderia ser a base de uma nova empresa. Pouco tempo depois, em 18 de julho de 1968, Moore e Noyce fundaram a empresa **NM Electronics** (**N** de Noyce e **M** de Moore). Nesse dia contrataram para diretor de engenharia [Andrew Grove](#) (1927 – 1990) para implementar a nova empresa. Pouco tempo depois, mudaram o nome para **Integrated Electronics**, cujo logo seria INTEL, mas tiveram de pagar 15000 dólares a uma empresa que já tinham registado esse nome.



Grove, Noyce e Moore  
**Fundadores da INTEL.**

Em abril de 1969, a INTEL comercializava o seu primeiro circuito integrado, uma memória de 64 bit com tecnologia inovadora (transistores bipolares [Schottky](#)). Em julho de 1969, a INTEL desenvolveu a sua primeira memória estática i1101 com 256 bit, já com a nova tecnologia MOSFET de canal P. Foi da resolução de problemas evidenciados por este circuito que se chegou ao 3º produto INTEL – a memória EPROM i1702 e, alguns meses depois, ao primeiro microprocessador monolítico, o INTEL i4004.

## Do TPF às primeiras memórias ROM e EPROM

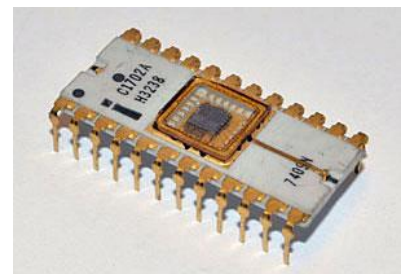
No outono de 1969, a empresa Intel tinha problemas de fiabilidade com as primeiras memórias estáticas MOSP, SRAM designadas por i1101, ([ver aqui as primeiras memórias monolíticas](#)).

Estes problemas foram investigados pelo israelita e estadunidense [Dov Frohman](#)<sup>5</sup> (1939 - ...) que acabou por descobrir a causa deste problema. Os isolantes de dióxido de silício do chip i1101 absorviam eletrões dos componentes metálicos e ficavam polarizados com carga elétrica, o que interferia no funcionamento do dispositivo MOSFET pois ficava permanentemente num dado estado de condução.

## EPROM, a extraordinária descoberta de Frohman

Depois de descobrir os defeitos provocados pelas cargas “estranhas” armazenadas na memória INTEL i1101, Frohman pensou que se induzisse carga elétrica, no óxido de silício, poderia criar um circuito programável eletricamente, que poderia reter a carga elétrica sem a necessidade de a memória ser alimentada continuamente.

Frohman, também verificou que a luz ultravioleta, UV, fazia desaparecer as cargas armazenadas no óxido de silício. Foi assim que chegou à célula de memória do TPF, PROM, usada para construir as memórias i1601 e à EPROM i1702 numa arquitetura de 256 x 8 bit, Fig. 9.



**Fig. 9 - Intel 1702, fevereiro de 1971**

A diferença entre estes dois circuitos é que a i1702 tem uma janela de quartzo que permite que a exposição prolongada a uma luz ultravioleta, com o comprimento de onda de 200 nm, apague a memória para poder voltar a ser programada.

Com a tecnologia de 10 um, disponível, em 1971, cada transistor ocupava uma área quadrada com cerca de 45 um de lado e a espessura do óxido fino era de cerca de 50 a 100 nm. Devido a esta grande espessura era impossível usar o tunelamento quântico microscópico para induzir cargas na PF do transistor.

<sup>5</sup> Frohman foi o fundador da Intel Israel, a primeira fora dos EUA.





Frohman, usou o método de injeção de portadores quentes, HCI, acelerando os eletrões com uma tensão de -48 V, que era aplicada na porta de controlo do transístor, relativamente ao substrato.

### A brilhante primeira demonstração da EPROM

Frohman, apresentou as capacidades da EPROM numa demonstração ao vivo, pela primeira vez, na conferência “*International Solid State Circuits, ISSCC 1971*”, utilizando os bits de informação programados na memória Intel 1601, com janela de quartzo. Frohman, usou um controlador para ler o logo "Intel" e mostrá-lo num ecrã de um osciloscópio Tektronix, Fig. 10.

De seguida, Frohman iluminou a memória 1702 (a 1601, com janela de quartzo) com luz UV bastante forte.

Segundo o relato de Gordon Moore, Dov Frohman projetou um filme que mostrava o padrão de bits do logo INTEL nas células de memória EPROM.

*Dov projected a film that displayed the bit pattern in the EPROM memory cells. As the cells were exposed to ultraviolet light, the bits dropped out one-by-one until all that was left was the familiar Intel logo. ... The bits fell, and when the final one disappeared, the entire audience broke into applause. Dov's paper was voted the best at the conference.*

Frohman patenteou a ideia com o nº US 3744036. Em 1972, a INTEL começou a comercializar este circuito revolucionário de que o Museu Faraday tem dois exemplares<sup>6</sup>. O preço deste circuito inovador era muito elevado (equivalente a cerca de 4000 €, em 2026)



Fig.10 - Dov Frohman na 1ª demo de uma memória EPROM (1971).

### O impacto da EPROM na indústria eletrónica

O principal impacto da i1702 não foi inicialmente percecionado pois julgava-se que o interesse maior seria para os engenheiros que desenvolviam programas em computadores dedicados e terem a possibilidade de mudar o código rapidamente programando de novo o circuito.

Com o aparecimento do 1º microprocessador monolítico o INTEL 4004, desenvolvido por [Federico Faggin](#) (1941 - ...), [Ted Hoff](#) (1937 - ...) e [Stan Mazor](#) (1941 - ...), em novembro de 1971, alguns meses depois da memória i1702, o panorama mudou, pois passou a ser vulgarizada a possibilidade de qualquer pessoa fazer o seu computador e fazer a sua reprogramação rápida.

Pode, também, ter acesso ao [documento do Prof. Rui Duarte](#), relativo aos 50 anos do microprocessador e do estabelecimento da célebre lei de Moore.

Gordon Moore, Robert Noyce e Andrew Grove, consideraram a EPROM "tão importante para o desenvolvimento da indústria de microcomputadores quanto foi o próprio [microprocessador](#)".

O ano de 1971 foi [o primeiro em que a Intel teve lucros](#). A EPROM, até ao fim da década de 80, foi o produto mais lucrativo da Intel.

No Museu Faraday do IST pode observar a pastilha de um i4004, num relógio comemorativo dos 25 anos do microprocessado i4004, Fig. 11, que foi oferecido em 1996 aos engenheiros seniores da INTEL.

<sup>6</sup> - Os dois circuitos, em 1972, custaram o valor do ordenado de um assistente eventual (cerca de 6000 escudos, aproximadamente 4100 € de 2026).



Fig. 11- Museu Faraday do IST  
Relógio comemorativo dos 25 anos do i4004.

### Programação das EPROMS

A EPROM i1702 precisava de várias tensões aplicadas aos terminais e algumas regras de sequenciação das operações e das suas temporizações resumidas na Fig. 12.

Alguns terminais da i1702 eram usadas na leitura da informação contida no circuito (valores indicados a preto) mas, na fase de gravação, poderiam ser submetidos a tensões diferente (valores de tensão indicados a vermelho).

| Pino do Chip     | Nome Original                                | Tensão no Modo Leitura | Tensão / Estado no Modo Programação                                                                                           |
|------------------|----------------------------------------------|------------------------|-------------------------------------------------------------------------------------------------------------------------------|
| Pino 24          | $V_{CC}$                                     | +5 V                   | +5 V (Mantém-se fixo) [Intel Silicon Gate MOS 1602A/1702A-S714]                                                               |
| Pino 12          | $V_{DD}$                                     | -9 V                   | GND (0 V) [Intel Silicon Gate MOS 1602A/1702A-S714]                                                                           |
| Pino 13          | $\overline{CS}$ (Program)                    | GND / +5 V             | Pulso de Programação de 0 V a -48 V [Intel Silicon Gate MOS 1602A/1702A-S714]                                                 |
| Pino 14          | Desconectado                                 | N/C                    | $V_{BB} = +12$ V (Substrato - Corrente limitada a 100mA max) [Intel Silicon Gate MOS 1602A/1702A-S714, Programming - Tronola] |
| Pino 15          | Desconectado                                 | N/C                    | $V_{GG}$ Pulsado (Alterna entre 0 V e -40 V / -48 V) [Intel Silicon Gate MOS 1602A/1702A-S714]                                |
| Pinos 1-3, 17-21 | Endereços (A <sub>n</sub> a A <sub>4</sub> ) | TTL (0 V ou 5 V)       | Lógica Invertida Pulsada (0 V para nível Alto / -48 V para nível Baixo)                                                       |
| Pinos 4-11       | Dados (D <sub>n</sub> a D <sub>4</sub> )     | Saída TTL              | Entrada Pulsada (Onde injeta os dados a gravar usando pulsos de -48 V)                                                        |

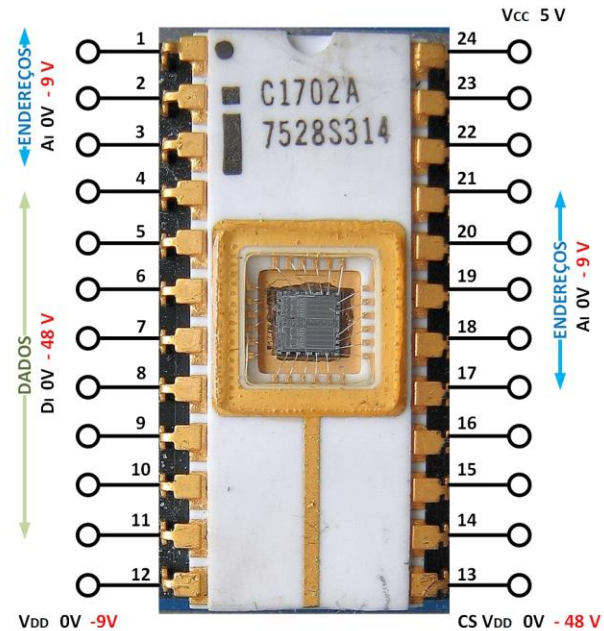


Fig. 11 – Informação básica sobre a i1702. Função dos terminais e tensões usadas.

Durante a década de 70 houve uma grande evolução nas memórias de TPF. Estas memórias necessitavam de várias tensões de alimentação, algumas positivas outras negativas. A passagem para substrato do tipo P permitiu fazer TPFs de canal N e obter uma melhoria considerável de desempenho, sendo o primeiro exemplo a memória 2708 de 1975, ver Tab.2, desenvolvida por [George Perlegos](#) (1950 - ...) da INTEL.

A memória 2716, introduzida em 1977, reduziu a tensão de leitura para uma única fonte de 5 V, mas a tensão de programação e as temporizações tinham de vir do exterior.



Entre 1976 e 1978, [Eli Harari](#)<sup>7</sup> (1945 - ...) na [Hughes Aircraft Company](#) apercebeu-se que se a espessura do óxido fino do TPF fosse reduzida para 10 nm seria possível usar o tunelamento de Fowler-Nordeim, quer para escrever, quer para apagar a célula de memória baseada no TPF. Harari submeteu, em fevereiro de 1977, a patente do processo que tinha desenvolvido e, em setembro de 1978, a patente foi concedida, o que permitiu à Hughes fazer uma demonstração oficial do processo.

Entre 1979 e 1981, Harari trabalhou na INTEL como chefe do grupo de desenvolvimento das memórias não voláteis de onde viria a sair a primeira memória de escrita e apagamento elétrico.

Em 1980, George Perlegos, na INTEL, desenvolveu a memória 2816, introduzindo pela primeira vez<sup>8</sup>, em produção, o método de programação e apagamento elétrico EEPROM (*Electrically Erasable Programmable Read-Only Memory*) por tunelamento quântico de Fowler-Nordheim, mas que precisava de 2 transístores por célula de memória. A tensão de programação de 21 V ainda tinha de ser aplicada do exterior, bem como as temporizações necessárias para a programação, mas esta tecnologia viria a revolucionar o mercado das memórias de TPF.

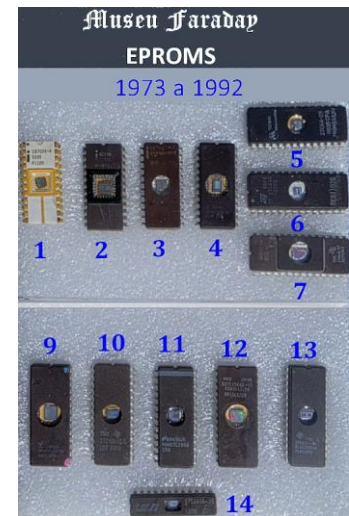
| Modelo     | Ano  | Tecnologia | Tensão de Leitura | Tensão de Gravação | Método de Apagamento       | Ciclos de Vida Estimados |
|------------|------|------------|-------------------|--------------------|----------------------------|--------------------------|
| Intel 1702 | 1971 | PMOS       | +5V, -9V          | -48V               | Luz UV (Janela de quartzo) | 10 - 20                  |
| Intel 2708 | 1975 | NMOS       | +5V, -5V, +12V    | +26V               | Luz UV (Janela de quartzo) | ~100                     |
| Intel 2716 | 1977 | NMOS       | +5V apenas        | +25V               | Luz UV (Janela de quartzo) | 100 - 1.000              |
| Intel 2816 | 1980 | EEPROM     | +5V apenas        | +21V (Elétrica)    | Totalmente Elétrico        | 10.000 a 100.000         |

**Tab. 2 . -Evolução das EPROM's INTEL**

Em 1981, Perlegos e outros ex-funcionários da INTEL formaram a empresa [SEEQ Technology](#)<sup>9</sup>, onde viriam a desenvolver a primeira memória programável e apagável eletricamente, a SEEQ 5213, sem a tensão de alimentação externa de programação ( este circuito gerava internamente a elevada tensão de programação com um circuito elevador de tensão eletrónico do tipo *Charge Pump*).

Em 1983, a Intel apresentou a revolucionária memória EEPROM 2816A, também usando o princípio da *charge pump* e que libertava os barramentos externos de endereço e de dados, continuando o processo de gravação e apagamento a correr internamente, em segundo plano, sem que o processador a que estava ligada percebesse. Estava, assim, aberto o caminho para as novas gerações de memórias de TPF do tipo *Flash*.

No Museu Faraday do IST pode encontrar um expositor, Fig. 12, com as memórias EPROM mais notáveis, que foram desenvolvidas entre 1973 e 1992. Estas memórias ilustram bem a evolução da tecnologia pois com pastilhas de silício do mesmo tamanho a capacidade de armazenamento creceu de 2 kB a 512 kB em 20 anos.



**Fig. 12 - Da EPROM de 2 kB a 512 kB.**

### O apagamento das EPROMs

Frohmman experimentou deixar as memórias expostas diretamente ao Sol durante vários dias e elas mantiveram a informação sem erros.

A exposição prolongada à luz ultravioleta natural do Sol pode apagar gradualmente os dados armazenados, por isso, os *chips* tinham uma janela de quartzo transparente protegida por um autocolante opaco.

<sup>7</sup> Eli Harari tem um trabalho excecional no desenvolvimento de memórias não voláteis de que detém mais de 1000 patentes.

<sup>8</sup> - Este método tinha sido testado entre 1976 e 1978 pela empresa Hughes Microelectronics.

<sup>9</sup> - A SEEQ viria a ser adquirida pela empresa LSI Logic.





Froham verificou, também, que a luz UV com 200 nm de comprimento de onda era muito eficaz a apagar as memórias EPROM.

Nos anos 70 e 80, as EPROMs foram usadas para guardar o código dos microprocessadores. Eram escritas através de programas de computador e máquinas de programação/cópia que proporcionavam os sinais e tensões de programação adequados, ver exemplo na Fig.13.



Fig. 13- Programador/copiador de EPROM e apagador.

As correções dos programas podiam ser feitas através do apagamento da EPROM numa fonte de luz UV e duma reprogramação corrigida.

### O efeito seminal da memória EPROM na renovação empresarial

A exploração das propriedades dos transístores de porta flutuante teve um efeito extraordinário na criação de novas empresas, tanto no domínio digital como analógico.

Em 1984, George Perlegos fundou a empresa [Atmel](#), *Advanced Technology for Memory and Logic*, que viria a desenvolver o primeiro microprocessador com memória programável, não volátil, integrada do tipo EEPROM. Esta empresa criou a linha de processadores AVR com memória flash integrada, A Atmel viria a ser comprada, em 2016, pela [Microchip Technology](#).

Em 1984, o engenheiro [Fujio Masuoka](#) (1943 - ...), da empresa japonesa Toshiba, apresentou o primeiro circuito de memória Flash do mundo, do tipo NOR, mas os responsáveis da Toshiba não viram interesse comercial na invenção; já a empresa INTEL viu aqui uma grande oportunidade comercial.

Em 1985, a INTEL inventora da DRAM abandona [este mercado que tinha criado](#), o que a salvou da falência. Começou a apostar forte no mercado dos microprocessadores, que também tinha criado<sup>10</sup> Toshiba, Hitachi, NEC e Fujitsu vendiam as DRAM abaixo do preço de custo americano. A participação da Intel no mercado global de memórias caíra de 83%, em 1974, para 1,3% em 1984, e, em 1985, o lucro por ação foi de um centavo de dólar. No ano seguinte, o prejuízo chegou a US\$ 107 milhões.

Em 1985, na INTEL, [Stefan Lai](#) co-inventa a tecnologia das memórias Flash ETOX (EPROM *Tunnel Oxide*). Junta-se a célula de um único transístor da EPROM, com o apagamento/programação da EEPROM (2 transístores) numa célula com um único transístor modificado com a adição de duas portas sobrepostas. Esta memória usava a injeção de portadores quentes HCI na programação, mas para o apagamento usava o tunelamento de Fowler-Lordeim.

A estrutura da memória ETOX era do tipo NOR, que permitia ler bit a bit de forma extremamente rápida, mas o apagamento da informação era feito por grandes blocos e não se podia apagar individualmente uma célula. Stefan Lai é, atualmente, um dos especialistas da empresa [Zeno Semiconductor Inc.](#) A tecnologia Intel Etox teve bastante sucesso como memória de programa de telefones celulares.

Em 1985, [Yoshishige Kitamura](#), da empresa NEC, registou a primeira patente para a tecnologia de células multiníveis (MLC) para armazenamento em EPROM. Numa única célula de memória, cada transístor de porta flutuante armazena-se mais do que um bit de dados, correspondente a vários níveis de carga armazenada.

Em 1987, [Fujio Masuoka](#) (1943 - ...), da empresa japonesa Toshiba [apresentou o conceito da memória Flash NAND](#), pela primeira vez, que permitia fazer circuitos muito mais compactos do que a sua tecnologia NOR. O termo Flash viria a ser introduzido por um colega de Fujio Masuoka que notou a extrema rapidez com que

<sup>10</sup> Ver o livro de Andy Grove "Only the Paranoid Survive".



se podia apagar a memória NAND, quase como a rapidíssima duração do flash de luz de uma máquina fotográfica. A tecnologia NAND da Toshiba começou a concorrer com a tecnologia NOR da INTEL.

### As memórias permanentes tipo NOR e NAND com TPF

Os dispositivos de armazenamento de carga (células) podem ser organizados em palavras digitais com um certo número de bits cada. Mas na implantação física há dois modos fundamentais: NAND e NOR, ver Fig. 14.

No modo NOR cada célula da memória pode ser endereçada individualmente, para escrita ou para leitura, atuando sobre a linha bit da palavra e selecionando a palavra cujo bit se quer ler/ escrever. Ex: ler bit 3 da palavra 4.

Neste exemplo de 8 células de memória colocadas lado a lado, há 17 ligações de linhas elétricas, que ocupam muito espaço na implementação.

O modo NOR permite ter acesso aleatório individual a qualquer célula, ou conjunto de 8 (Byte) e lêr dados de forma muito rápida sem o código passar para uma memória RAM de acesso temporário. Este tipo de memórias é excelente para guardar *firmware*, BIOS e sistemas de arranque de dispositivos computacionais.

No modo NAND de implementação as células estão ligadas em série e não há acesso individual ao conteúdo de uma dada célula. Isto permite ter muito mais células no mesmo espaço. Na NAND a escrita grava e apaga por blocos de dados de forma muito rápida. É a organização ideal para *pen drives* cartões de memória e “discos” de estado sólido – *Solid State Disc*, SSD.

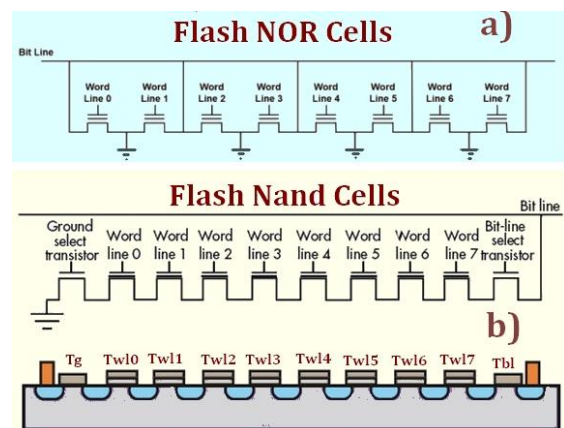


Fig. 14 -Memória organizada no modo NOR e no modo

### A exploração comercial das memórias NOR e NAND

Em 1987, a empresa [General Instrument](#) separou a sua divisão de microeletrônica, que foi adquirida em 1989, por um grupo de investidores (Sequoia Capital). Em 1990, a nova empresa começou a ser liderada por [Steve Sanghi](#) que também tinha saído da General Instrument. Steve valorizou muito a empresa que passou a chamar-se [Microchip Technology](#).

Em 1988, Eli Harari, [Sanjay Mehrotra](#) (1958 - ...) e [Jack Yuan](#), criam a empresa [SanDisk](#), inicialmente sob o nome de SunDisk e a empresa progrediu muito rapidamente.

Em 1993, a [Atmel](#) produziu o primeiro microcomputador com memória flash integrada, do tipo NOR, o AT89C51, baseado na arquitetura do núcleo INTEL 8051.

Nos anos 90 a Toshiba, aliou-se à Samsung e à Sandisk para desenvolverem a tecnologia NAND Flash e combaterem a tecnologia NOR Flash da INTEL.

Em 1996, a Sandisk produziu a primeira memória flash MLC, mas em 2015, depois do enorme sucesso, viria a ser adquirida pela empresa [Western Digital](#) por 19 bilhões de dólares EUA.

Em 1997, a Intel comercializou a série de memórias [Flash Intel Strata Flash](#), uma tecnologia MLC com 4 níveis de carga em cada transístor, sendo a primeira memória NOR MLC de grande produção (2 bits /célula).

Com o aparecimento dos dispositivos de armazenamento de grandes ficheiros, como imagem e vídeo, dispositivos de som codificado no formato mp3 e *pens* USB, a tecnologia Intel NOR Flash deixou de ser competitiva e focou-se em mercados de nicho. A empresa [Numonix](#), formada em 2008, comercializou as





memórias NOR da Intel, as memórias Flash NAND da [ST Microelectronics](#) e a participação dos investidores [Francisco Partners](#), foi comprada pela Micron Technology em 2010 por 1,27 mil milhões de dólares.

### As memórias Flash em formato de cartões

Em 1994, a SanDisk introduziu o formato CompactFlash (CF), com memória Flash do tipo NAND SLC. A CompactFlash, Fig. 15 A), simulava um disco rígido comum (Interface ATA/IDE:) O controlador estava integrado no cartão permitia que ele se comportasse exatamente como uma unidade de armazenamento de disco rígido de um computador e faz a interface direta com um dispositivo, tipo câmara de filmar ou fotográfica.

O cartão tinha 50 pinos do tipo fêmea para ligar aos pinos macho do dispositivo. Estes pinos são um subconjunto dos pinos da interface de 68 pinos PCMCIA. Para fazer a interface com as fichas USB foram feitos adaptadores, Fig. 16, para ligar aos portos USB dos computadores.

Em 1995, a Toshiba, lançou o formato SmartMedia (originalmente chamado de Solid State Floppy Disk Card ou SSFDC), Fig. 15 B), com memória flash do tipo SLC NAND), mas não tinha controlador interno função que teria de ser do dispositivo computacional, o que dificultava a sua universalidade, tendo o formato acabado por morrer em 2003.

Em 1998, a Sony lança o Memory Stick, Fig. 15 C), com memória flash do tipo SLC com a capacidade inicial de 4 MB e 8 MB. Este dispositivo já tinha um controlador interno.

Em 1999, as empresas Panasonic, SanDisk e a Toshiba, uniram-se e apresentaram o formato Secure Digital (SD) que tinha controlador interno e um sistema de encriptação de direitos digitais. Este cartão acabou por conquistar o mercado. A arquitetura do SD permitiu criar novas gerações retro compatíveis: SDHC (até 32 GB), SDXC (até 2 TB) e o recente SDUC (até 128 TB).

### O novo negócio dos Disc On Chip e PenDrive

A criação das memórias Flash de tamanho físico muito pequeno foi logo visto como um negócio muito apetecível pois poderia substituir o sistema clássico de disquetes usado para transportar informação digital. Houve uma grande luta de patentes que acabou por expirar em 2020.

Em março de 1993, [Amir Ban](#), tinha submetido a patente US 5404485, que foi concedida em abril de 1995, designada por *Flash File System* que seria renomeada para *Flash Transfer Layer*, FTL. A FTL é um método implementável (em *hardware*, *firmware* ou *software*) para controlar e gerir o acesso a uma memória *flash*, ultrapassando as suas restrições de funcionamento, de modo a que esta seja vista pelo sistema operativo de um computador como um dispositivo de armazenamento de dados no qual é possível ler e escrever dados em qualquer local da memória flash.

Inicialmente, Ban quando criou o sistema tinha em vista as memórias flash tipo NOR mas, posteriormente, adaptou os algoritmos às memórias flash tipo NAND.

Em 1994, [Ajay Bhatt](#) (1957 - ...) na IBM, começou a desenvolver a interface de computadores *Universal Serial Bus*, USB, a que se juntaram vários fabricantes e, em 1996, saiu a primeira norma USB 1.0.



Fig. 15 - Primeiros cartões de memória Flash.



Fig. 16 - Adaptador CompactFlash USB



Em 1995, a [M-Systems](#), fundada pelos israelitas [Dov Moran \(1955 - ...\)](#), [Amir Ban](#) e [Oron Ogdan](#), criou o *Disk On Chip*, DOC, Fig. 17. O DOC era um circuito de encapsulamento DIL, *Dual In Line*, de 32 pinos, que continha memória barata tipo *flash* NAND a que era acrescentado um pequeno circuito controlador baseado no FTL. O DOC era visto como uma memória permanente NOR convencional, mas muito mais barata. A memória flash era fornecida pela Samsung Eletronic.

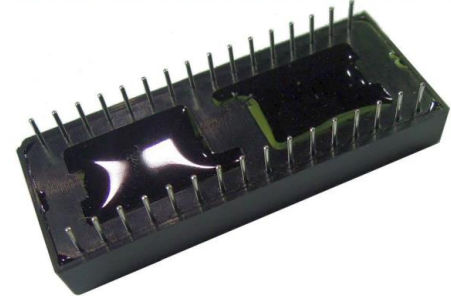


Fig. 17– DOC de 16 MB.

Ainda em 1995, os israelitas [Dov Moran](#), [Amir Ban](#) e [Oron Ogdan](#) patentearam, [US 6148354A](#), um equipamento que designaram como “*USB-based PC flash disk*”. Criaram a empresa [M-Systems Flash Disk Pioneers Ltd](#), que desenvolveu a arquitetura de armazenamento flash baseada na interface USB e registaram várias patentes do dispositivo. Nasceu, assim a primeira *Pen Drive*.

O sistema *Flash Transfer Layer* de Ban já incluía as funções de Recolha de lixo, *Garbage Collection*, Escrita fora do sítio, *Out-of-place writes*, Nivelamento do desgaste das células, *Wear Leveling*, Correção de erros, *Error Correction Codes* e outras funções que tornariam as futuras memórias flash resilientes e duradouras.

Em 1995, Shimon Shmueli da IBM criou um sistema parecido ao da M-Systems, que ligava a memória flash, do tipo SLC, diretamente à porta USB. A IBM não patenteou o resultado, mas fez uma parceria estratégica com a M-Systems que fabricaria os dispositivos e a IBM trataria da divulgação e venda através da sua gigante rede comercial

Em 1995, Pua Khein-Seng, CEO da [Phison Electronics Corporation](#), com sede em Taiwan, também afirmou ter inventado um dispositivo com as mesmas características do criado por Dov Moran em Israel.

Em 1996, a Intel introduziu uma memória Flash considerada revolucionária, focada na utilização em dispositivos móveis, designada por 28F016SV, de 16 Mbit realizada com tecnologia ETOX de 600 nm, mas que era capaz de suportar 1 milhão de ciclos de escrita e apagamento. Esta memória estava dotada de um sistema SV “*Smart Voltage*” com possibilidade de escolha de tensão de alimentação de 3,3 ou 5 V e de programação de 5 V ou 12 V. Era garantida como sendo mais fiável do que o disco rígido de 1”, que começava a ser usado nos telemóveis.



Fig.18 - Intel Smart Voltage Flash

Uma das primeiras unidades comercializadas pela Intel, Fig. 18, foi usada no sistema<sup>11</sup> “[On Road Dyno](#)” que faz parte do repositório do Museu Faraday.

Em 1999, a empresa chinesa [Netac](#) clamava que tinha inventado a Pen drive, mas a sua unidade só chegaria ao mercado em 2002, não havendo provas de ter sido nessa data.

Em 1999, o padrão de interface a associação de construtores [PMCIA](#) adotou o sistema de Ban, designado agora por *True Flash File System*, TFFS, que hoje em dia, é usado por todos os sistemas de armazenamento Flash.

<sup>11</sup> [On Road Dyno](#) “Banco de Ensaio Portátil para Veículos Motorizados” António Oliveira, José Jordão e Renato Caldeira, TFC 1996/97 Engenharia Eletrotécnica e Computadores do IST.



Em 2000, a empresa [Trek 2000 International](#), de Singapura, foi a primeira a vender uma unidade flash USB, chamada ThumbDrive, Fig. 19. A Trek acabou por ganhar a corrida comercial.

Em 2000, a M-Systems e a IBM começaram a comercializar o “DiskOnKey”, Fig. 20, que teve sucesso imediato. A IBM vendia o sistema com o nome “Memory Key”, primeiro no formato de 8 MB, depois 16 MB e 64 MB (em 2003), produzido pela M-Systems.

A IBM tinha inventado a disquete de 3,5” em 1981, mas com o aparecimento das *pen drives* e a parceria que fez com a M-Systems as disquetes começaram a não ter mercado e fechou a produção de disquetes em 2011.

Em 2006, a IBM acabou por sair deste mercado das *pen drives* e Dov Moran vendeu a empresa M-Systems à SanDisk por 1.6 bilhões de dólares.

### A EEPROM no domínio analógico

Sendo analógica a informação da carga armazenada nos transístores de porta flutuante, as memórias digitais feitas com esta tecnologia implicam ter uma entrada digital que é convertida para um valor analógico multinível que será armazenado no transístor e, no modo de leitura, este valor é convertido numa informação digital. Na realidade, as EEPROM são memórias analógicas adaptadas para uso digital, que é realmente o seu grande mercado.

Em 1978, [Raphael Klein](#), Richard Simko e outros ex-especialistas da INTEL formaram a empresa XICOR para explorar aplicações dos TPF. O primeiro produto foi a série de memórias NOVRAM (*Non-Volatile Static RAM*), que retinham dados caso houvesse uma falha repentina da tensão de alimentação. Depois, a XICOR especializou-se na realização de potenciômetros digitais usando os transístores de porta flutuante para reter a informação do valor de resistência do potenciômetro.

Os potenciômetros digitais, usualmente designados por *digipots*, são constituídos por agregados de resistências divisoras de tensão ligadas por MOSFETS e que podem ser controlados por interruptores manuais ou por comandos de um processador através das portas SPI ou I2C. Podem incluir memórias não voláteis do tipo EEPROM que retêm os valores desejados de atenuação de sinais analógicos e, usualmente, incluem um processador para fazer a calibração e a correção de erros do divisor resistivo, como, por exemplo, o potenciômetro AD5141BCPZ10, Fig. 21, da [Analog Devices](#)<sup>12</sup>

Este circuito armazena o erro da tolerância das resistências na sua memória EEPROM o que permite ao controlador interno corrigir a atenuação com uma precisão de 0,01%<sup>13</sup>.

Em 2004, a Xicor foi adquirida pela empresa estado-unidense Intersil Corporation que, por sua vez, em 2017, foi adquirida pela japonesa [Renesas Corporation](#) que atualmente explora o negócio dos potenciômetros digitais baseados no TPF.



Fig.19- Trek Thumb Drive (8 MB)



Fig. 20- IBM Disc On Key.

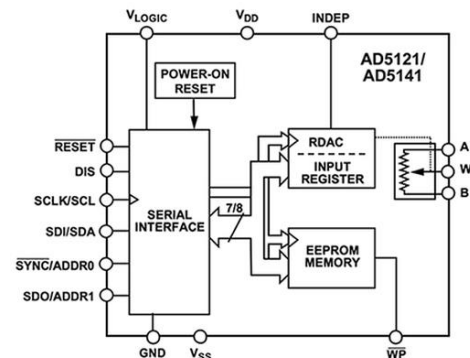


Fig. 21- Digipot AD5141BCPZ10.

<sup>12</sup> A empresa AD tem vindo a comprar várias empresas especializadas em circuitos analógicos e lidera este mercado.

<sup>13</sup> Essa precisão é cem vezes melhor do que a precisão de 1% dos *digipots* de maior precisão sem NVM.

[illegible]

O sucesso foi imediato com o chip [ID 1820](#) designado por ChipCorder, que permitia guardar diretamente 20 s de sinais de voz analógica, sem o uso de qualquer processador digital, Fig. 22.

A partir do ano 2000, começou a reaparecer muita investigação na tecnologia do transístor de porta flutuante para realizar sistemas de memória analógicos, quer para [gravação de sinais](#) quer para realizar circuitos de redes [neurais analógicas](#) muito mais rápidas que as redes digitais, bem como o seu uso em muitas outras aplicações.

Na “A gaiola eletrostática de Bits II” referiremos como a redução da dimensão dos transístores, e o desenvolvimento de células de armazenamento 3D V-NAND, com, 3, 4 e até 5 bit, por transístor, introduzem muitos erros nas memórias. Contudo, estes erros podem ser minimizados usando avançados controladores inteligentes associados à memória, códigos corretores de erro e o uso da inteligência artificial no controlo dos processos de gravação, apagamento e no conhecimento dos hábitos do utilizador. Estas técnicas permitem aumentar a fiabilidade destas memórias tanto em ambiente consumidor como no profissional.

Vídeo: - Gaiola Eletrostática de Bits.mp4