



Cosa serve per costruire una stampante 3D?

EOS Book

#23

Settembre 2015

Programmiamo

Raspberry Pi

con tutti i linguaggi



python

php

Shell

Bash



Java

LTE: scopriamo la telefonia mobile 4G

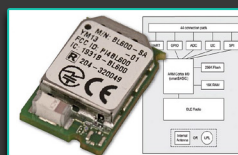
Antenne trasparenti: la tecnologia del mimetismo

Un portento della matematica: la trasformata di Laplace

Progetti Fai da te



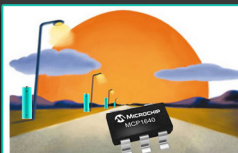
Gestiamo gli SMA con Arduino



Termometro Bluetooth a basso consumo



Il fornello solare



Crepuscolare a batteria



ELETRONICA OPEN SOURCE

Uno strumento di lavoro!



Erano i primi mesi del 2006, mi ricordo che stavo progettando FTPmicro (il primo webserver di 8 cm) e mentre leggevo una rivista di elettronica, una di quelle presenti allora in edicola, pensavo: "non è possibile, ancora lo stesso articolo, non ne posso più di tutorial sull'RS232!!".

Quella frustrazione nasceva proprio dal fatto che sulle riviste dell'epoca si leggevano sempre gli stessi articoli, ma non solo, in un contesto storico dove il web iniziava a prendere piede, questi articoli era facile trovarli liberamente online. Premetto che a quei tempi avevo tutte, dico tutte le riviste di elettronica disponibili, arretrati compresi. Mi ricordo che dalla mia assistente feci realizzare un foglio excel con tutti i numeri, titoli ed integrati utilizzati, in modo da poter risalire subito all'articolo relativo al progetto che stavo realizzando.

“ Ho da sempre considerato le riviste di elettronica uno strumento di lavoro!

Purtroppo però le mie aspettative furono deluse. Decisi allora di fare qualcosa, di realizzare io stesso degli articoli di elettronica in base alle mie esperienze e di coinvolgere amici progettisti elettronici nella condivisione della loro esperienza, dei loro progetti. Con il motto "share for life" nacque Your Electronics Open Source, quindi Elettronica Open Source.

Nello stesso periodo, qualcun altro forse ebbe un'idea simile dando vita a Firmware! Rivista che scoprii subito e da subito divenne la mia rivista di elettronica preferita e dopo poco, l'unica rivista italiana di elettronica che ho continuato a leggere.

Perché? Semplicemente perché è veramente

“ uno strumento di lavoro!

Con l'annessione di Firmware ad Elettronica Open Source, EOS-Book diventerà sempre più uno strumento di approfondimento a 360 gradi, con letture interessanti, sia per apprendere di nuove tecnologie, sia per realizzare progetti. Insomma una rivista per gli amanti della tecnologia, hobbisti e makers.

Il nostro impegno nel tempo sarà quindi di soddisfare le richieste e far accrescere il know-how sia

dei makers che dei professionisti.

Il primo numero della nuova gestione di Firmware sarà disponibile gratuitamente per tutti gli abbonati *Platinum*!

“ Avrete quindi a disposizione sia uno strumento di lavoro che uno strumento di divertimento che vi permetteranno di accrescere il vostro sapere in modo professionale e divertente allo stesso tempo!

P.S. Di RS232 ne abbiamo parlato anche su EOS ma in modo diverso

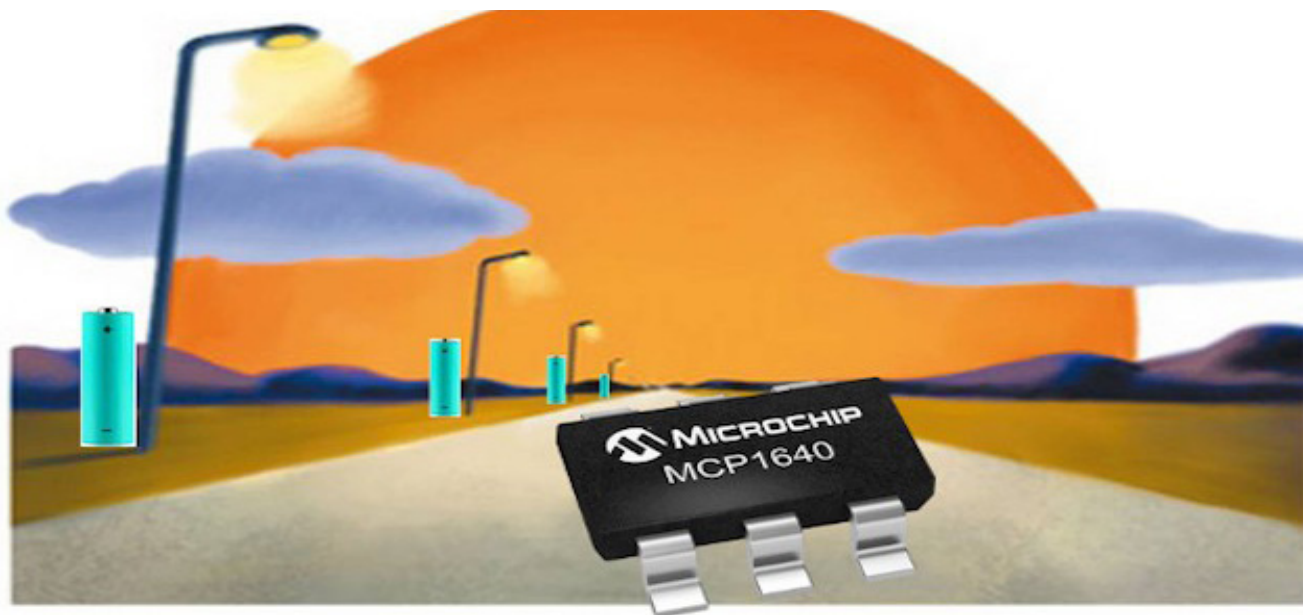
Emanuele Bonanni



Crepuscolare a batteria [progetto completo fai-da-te]



di Luca Giuliodori
1 Settembre 2015



Nel mondo dell'embedded l'utilizzo di accumulatori come fonte di energia è senz'altro un aspetto interessante e forse per alcune applicazioni è l'unica soluzione percorribile; molto spesso però si cerca di evitare l'uso delle batterie in quanto forniscono autonomia limitata e aumentano l'ingombro del dispositivo. Attraverso la progettazione e lo sviluppo hardware e firmware di un crepuscolare vedremo come, utilizzando l'MCP1640 della Microchip che può essere definito a tutti gli effetti un "minatore di cariche", questi limiti possono essere in parte superati; questo chip permette di ridurre il numero delle batterie pur mantenendo la tensione di alimentazione desiderata, facilita l'ottimizzazione dei consumi e soprattutto fornisce la tensione desiderata anche con batterie che solitamente vengono considerate scariche.

Nella progettazione di sistemi elettronici spesso ci si imbatte nell'esigenza di dover alimentare i dispositivi utilizzando appositi accumulatori per motivi legati sia alla praticità di installazione sia alla disponibilità della tensione

di rete. Se pensiamo ad esempio al mondo della sensoristica quanto appena detto risulta subito evidente e un esempio chiarificatore può essere quello dei sensori wireless per i sistemi di allarme: in molti edifici già esistenti e in cui si voglia

installare un impianto di allarme, non sempre è possibile raggiungere tutti i punti di accesso con dei sensori cablati; in questi casi sarà necessario adottare dei sensori wireless che dovranno lavorare senza cablaggio e quindi tramite l'utilizzo di batterie. Di applicazioni simili che utilizzano come fonte di energia gli accumulatori ce ne sono delle più svariate e presenti in diversi ambiti.

L'utilizzo di accumulatori al posto della tensione di rete è sicuramente una possibilità allettante in quanto, oltre ad evitare la realizzazione dei cablaggi per fornire l'alimentazione al dispositivo, semplifica il circuito da realizzare visto che la tensione è già in regime continuo e a valori relativamente bassi, caratteristica solitamente ricercata in quasi tutti i dispositivi elettronici e ottenuta tramite apposita circuiteria. L'entusiasmo per l'utilizzo degli accumulatori viene però subito attenuato dall'annoso problema della scarica delle batterie, in quanto implica inevitabilmente una manutenzione periodica da parte dell'utente o dell'installatore e quindi un maggiore costo di manutenzione rispetto ai dispositivi alimentati a tensione di rete. Questo "neo" non può essere evitato in quanto è insito nella natura stessa degli accumulatori può però in qualche modo essere limitato o ridotto: una delle soluzioni più adottate è senz'altro quella dell'**Energy Harvesting**, tecnica con cui si cerca in qualche modo di "raccolgere" tutta l'energia che il dispositivo ha a disposizione dall'ambiente esterno per poi utilizzarla come fonte di alimentazione; nel nostro caso non ci occuperemo di Energy Harvesting ma vedremo come sfruttare al meglio l'energia delle batterie senza doverle buttare non appena la tensione scende di qualche volt.

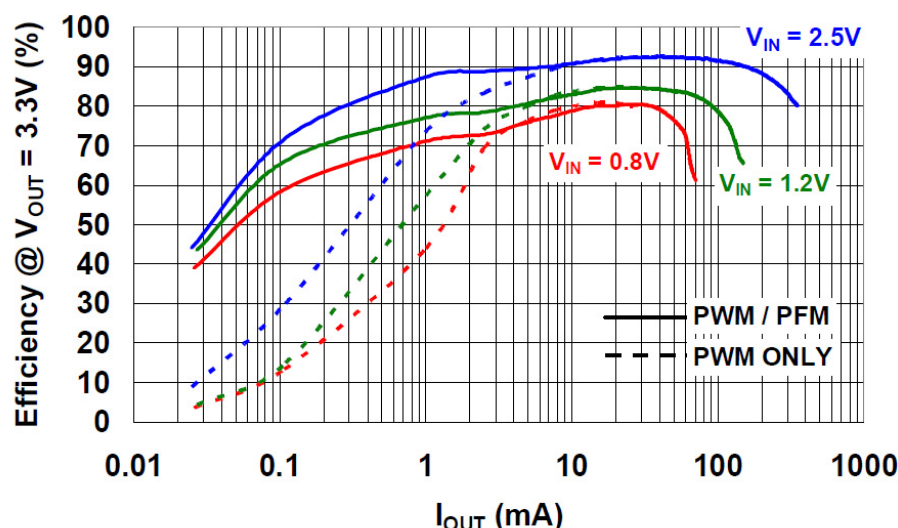
MCP1640

L'MCP1640 è un convertitore step-up DC-DC compatto che permette di ottenere, con un'efficienza che può arrivare fino al 96%, una tensione di uscita regolabile da 2,0V a 5,5V e con una corrente massima di 300mA, a partire da una tensione di ingresso variabile da 0,35V a 5,5V. Avere un range così ampio per la tensione di ingresso è una caratteristica molto interessante che si traduce in due grandi vantaggi:

- Il dispositivo continua a funzionare anche con batterie che solitamente vengono considerate scariche.
- Il numero delle batterie può essere ridotto.

Consideriamo ad esempio una generica applicazione che lavora con una tensione minima di 2V. Senza l'uso dell'MCP1640 questa dovrà essere alimentata ad esempio da una serie di minimo due batterie di tipo AAA da 1,5V; se questa serie scende al di sotto dei 2V (supponiamo 1V a batteria) il dispositivo smetterà di funzionare; utilizzando invece l'MCP1640 non solo sarà possibile adottare anche solo una batteria da 1,5V (a vantaggio dello spazio occupato ma a discapito ovviamente della durata), ma la batteria potrà essere sfruttata fino a 0,35V anziché 1V.

Una nota va fatta anche sull'efficienza di questo convertitore, infatti questa dipenderebbe fortemente dalla combinazione *tensione di ingresso - tensione di uscita - corrente di uscita* ma grazie ad un'opportuna **alternanza di due tecniche di controllo del convertitore (la PWM e la PFM)** si riesce a far mantenere il valore dell'efficienza a valori sufficientemente alti per un ampio range di tensioni e correnti, come mostrato dalla seguente figura:



La scelta di una tecnica di controllo piuttosto che di un'altra è operata di volta in volta in modo da massimizzare l'efficienza.

Questo chip per poter eseguire la conversione necessita ovviamente di energia la quale verrà inevitabilmente presa dalla batteria posta in ingresso e se questo consumo "passivo" risulta superiore o paragonabile al consumo dell'intero

dispositivo, influirà significativamente sulla durata della batteria. Per limitare questo effetto è stato previsto **un pin di *Enable*** che permette di mandare il chip in una sorta di stato di sleep in cui i consumi vengono drasticamente ridotti e portati al di sotto del micro ampere. Ovviamente quando

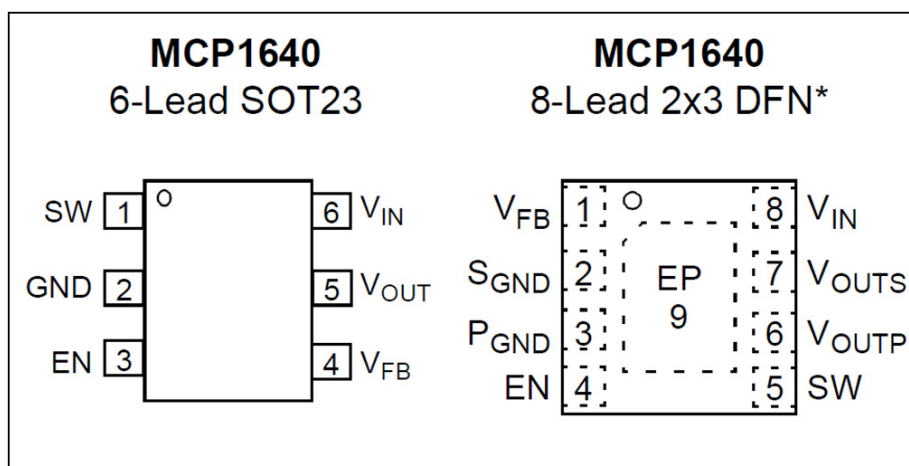
l'MCP1640 viene disabilitato il circuito ad esso connesso deve ridurre i consumi in quanto l'energia per mantenere il circuito in vita sarà quella contenuta nelle capacità di alimentazione; nel progetto che accompagna questo articolo verrà illustrato con maggiore dettaglio quanto appena

descritto.

Una cosa di cui fare molta attenzione è che, se viene sfruttata la funzione di *Enable*, il range di tensione di ingresso si riduce partendo non più dal valore 0,35V ma dal valore 0,65V, in quanto il convertitore è in grado di riabilitarsi solo se la tensione di ingresso è maggiore o uguale a 0,65V; ovviamente se il

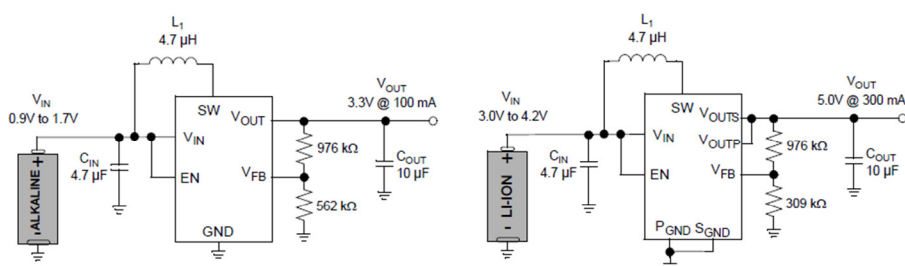
l'MCP1640 non viene disabilitato il range della tensione di ingresso rimane quello visto precedentemente.

L'MCP1640 è realizzato all'interno sia del package SOT23 a 6 pin sia del package DFN a 8 pin (9 se si considera anche il pin di dissipazione termica):



MCP1640 2x3DFN	MCP1640 SOT23	Simbolo	Descrizione
1	4	V_{FB}	Feedback di tensione
2	-	S_{GND}	Segnale di massa
3	-	P_{GND}	Segnale di power
4	3	EN	Enable
5	1	SW	Collegamento induttore
6	-	V_{OUTP}	Tensione di uscita
7	-	V_{OUTS}	Sens della tensione di uscita
8	6	V_{IN}	Tensione di ingresso
9	-	EP	Dissipazione termica
-	2	GND	Massa
-	5	V_{OUT}	Tensione di uscita

I pin tra i due package sono simili e per capirne il loro utilizzo una tipica applicazione sarà sicuramente chiarificatrice, come riportato nel manuale del componente e mostrato nella figura seguente:



Come possiamo vedere i **componenti che corredano l'MCP1640 sono veramente pochi e la maggior parte sono definiti da manuale:**

- Capacità di ingresso C_{IN} : il valore tipico per questa capacità è di $4,7\mu F$ come indicato nella figura, tuttavia il suo valore può essere aumentato in modo stabilizzare la tensione di ingresso.

- Capacità di ingresso C_{OUT} : per questa capacità il valore tipico è di $10\mu F$ ma possono essere utilizzati valori maggiori fino ad un massimo di $100\mu F$.
- Induttore L_1 : questo induttore può assumere un valore che va $2,2\mu H$ a $10\mu H$ anche se è raccomandato il valore $4,7\mu H$.

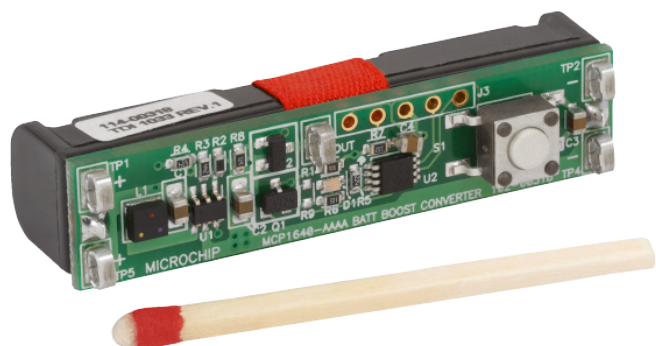
Resistenze poste in uscita: i valori di queste due resistenze permettono di selezionare la **tensione di uscita desiderata attraverso la formula $R_{TOP} = R_{BOT} (V_{OUT}/V_{FB} - 1)$** , dove R_{TOP} è la resistenza posta tra V_{OUT} o V_{OUTS} e V_{FB} , R_{BOT} è la resistenza posta tra V_{FB} e massa, V_{OUT} è la tensione desiderata e V_{FB} è la tensione di controllo che deve essere posta sempre al valore $1,21V$.

Per maggiori approfondimenti sul funzionamento dell'MCP1640 si rimanda alla lettura del manuale al seguente link: [MCP1640 Data Sheet](#).

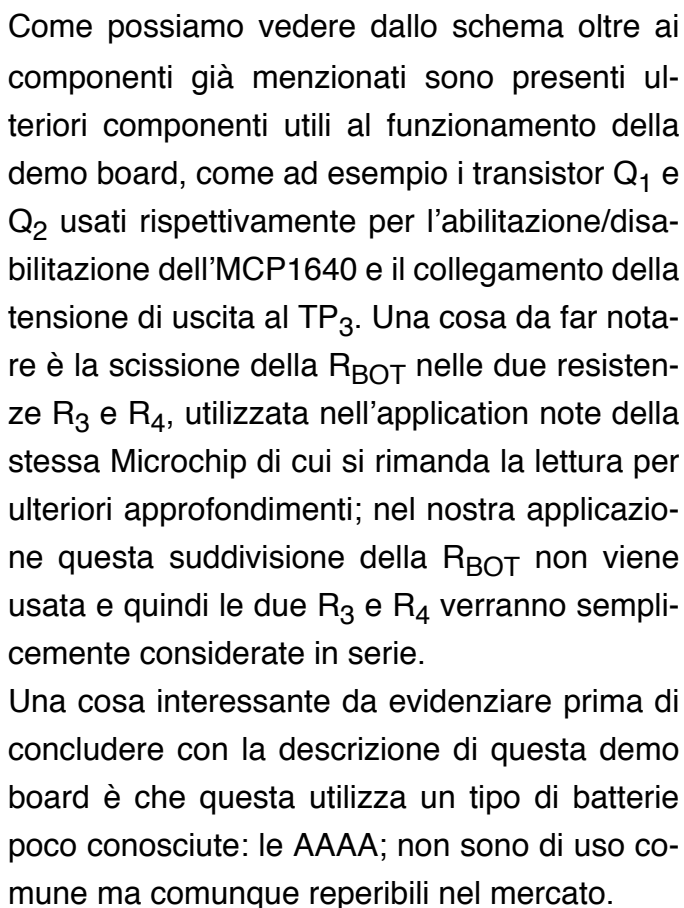
micro, come mostrato dallo schema elettrico riportato di seguito:

MCP1640 SINGLE QUAD-A BATTERY BOOST CONVERTER:

L'utilizzo di questo chip risulta molto interessante ma i package con cui viene fornito l'MCP1640 possono diventare un ostacolo se si intendono realizzare prototipi o dispositivi da hobbista. In questo caso la Microchip ci viene incontro fornendoci delle schede di sviluppo che rendono il dispositivo più accessibile e quindi utilizzabile; nel nostro caso si è utilizzata la MCP1640 Single Quad-A Battery Boost Converter:



Oltre alle dimensioni ridotte la cosa che rende veramente interessante questa scheda di sviluppo è che, oltre ad avere a bordo l'MCP1640 e tutti i componenti necessari a fornire una tensione di uscita pari a 3,3V, monta a bordo anche un led, uno switch, il **PIC12F617** con il software d'esempio dell'appliction note [AN1337](#) e il connettore J₃ per la programmazione e il debug del



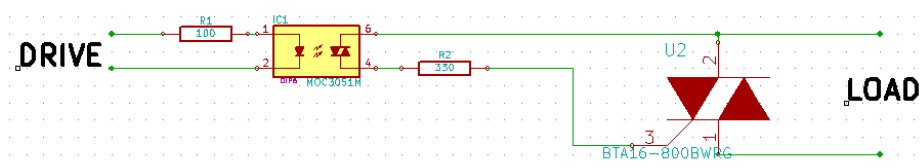
DESCRIZIONE GENERALE

CREPUSCOLARE

Da come si intuisce dallo schema la parte di logica e di gestione della batteria saranno implementate grazie all'utilizzo della demo board vista precedentemente. Chiaramente questo progetto non ha la pretesa di essere un crepuscolare da poter industrializzare ma è stato sviluppato a scopo illustrativo, per meglio comprendere il funzionamento del convertitore MCP1640.

HARDWARE

Come già accennato precedentemente la parte di logica e di gestione della batteria è affidata alla demo board MCP1640 Single Quad-A Battery Boost Converter della Microchip, mentre invece per quanto riguarda il sensore **si è scelto di adottare una semplice foto resistenza** e per la parte di potenza il **TRIAC BTA16-800BW** pilotato dal **MOC3051** come mostrato nella figura di seguito:



Il microcontrollore a bordo della demo board (il PIC12F617) dovrà quindi pilotare, tramite un pin configurato come uscita, il MOC per attivare il TRIAC e quindi il carico e leggere, tramite un pin configurato come ingresso e connesso all'ADC, il valore di tensione ai capi del foto resistore per scegliere se attivare o meno l'uscita che pilota il MOC. Con la demo board adottata **l'unico punto di accesso per far interagire il microcontrollore con il mondo esterno è il connettore J₃**, lo stesso usato per la programmazione e debug e in cui sono presenti i seguenti pin:

1. GP₃ (V_{pp})
2. V_{OUT}

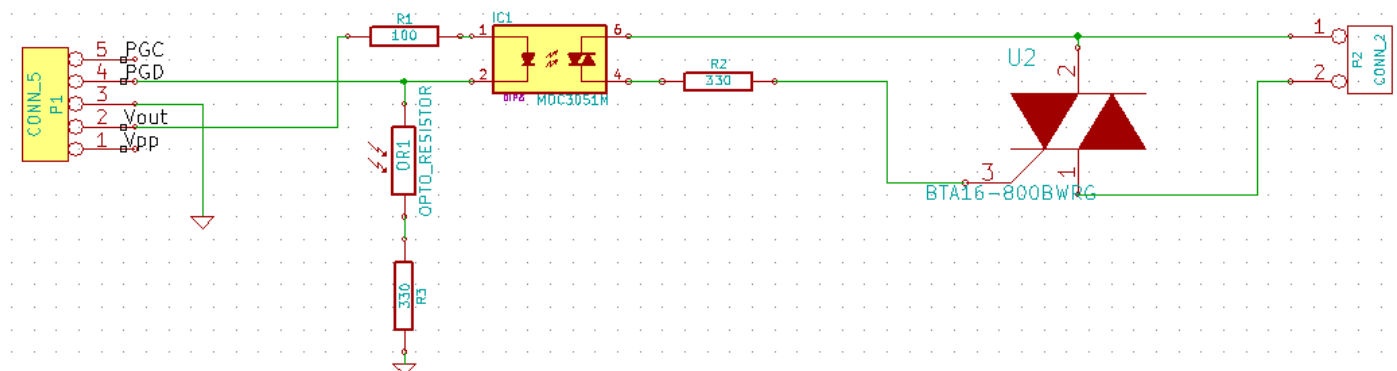
3. GND
4. GP₀ (PGD)
5. GP₁ (PGC)

Di questi 5 pin gli unici con cui il microcontrollore può interagire con il MOC e il foto resistore sono il GP₃, il GP₀ e il GP₁. Il primo può essere solamente configurato come ingresso e non è possibile utilizzarlo per la conversione A/D, caratteristiche che lo rendono inutilizzabile per i nostri scopi; il GP₁ può essere usato sia come ingresso sia come uscita ed è possibile usarlo per la conversione A/D ma siccome non è possibile alterare la tensione ai suoi capi in quanto connesso al partitore che seleziona la tensione di uscita dell'MCP1640 (fare riferimento allo schema elettrico già visto), anche questo pin risulta inutilizzabile; l'unico pin rimasto disponibili

le e che può essere configurato sia per il pilotaggio del MOC sia per leggere la tensione ai capi del foto resistore, senza compromettere il malfunzionamento dell'intero sistema, è il GP₀.

Ora però si ha il problema che con solo questo pin si devono svolgere due funzioni completamente differenti: pilotaggio MOC (che richiede di configurare il pin come uscita) e lettura tensione tramite ADC (che richiede di configurare il pin come ingresso e di usarlo per l'ADC). Per risolvere questo problema si è pensato di connettere al pin GP₀ il MOC in modo che questo si attivi solo quando il pin viene settato al valore logico 0 e, verso GND, la serie composta da un foto resistore e una resistenza da 330Ω, come mostrato dallo schema seguente:

Ora però si ha il problema che con solo questo pin si devono svolgere due funzioni completamente differenti: pilotaggio MOC (che richiede di configurare il pin come uscita) e lettura tensione tramite ADC (che richiede di configurare il pin come ingresso e di usarlo per l'ADC). Per risolvere questo problema si è pensato di connettere al pin GP₀ il MOC in modo che questo si attivi solo quando il pin viene settato al valore logico 0 e, verso GND, la serie composta da un foto resistore e una resistenza da 330Ω, come mostrato dallo schema seguente:

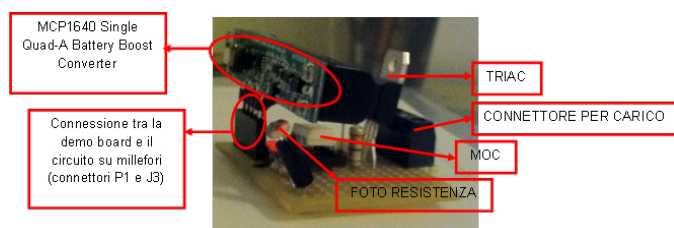


Ad onor del vero occorre specificare che la lettura della tensione ai capi della serie foto resistore - resistenza da 330Ω, **è possibile grazie alla presenza della resistenza di pull-up R8 prevista nella demo board.**

In questo modo quando si desidera attivare il MOC e quindi il TRIAC, sarà sufficiente configurare il GP₀ come uscita e settare il valore logico a 0, quando invece si desidera fare una lettura sul foto resistore, sarà sufficiente settare il pin come ingresso ed eseguire la lettura tramite ADC, senza causare malfunzionamenti dell'intero sistema o false attivazioni del MOC. La presenza della resistenza da 330Ω serve proprio a garantire che il MOC non si attivi nei casi in cui la foto resistenza assume valori prossimi a 0Ω. Chiaramente il connettore P₁ mostrato nello schema precedente è il corrispondente del connettore J₃ presente sulla demo board e con cui i due circuiti verranno connessi.

Una piccola nota va fatta sulla serie composta da R₈ – foto resistenza – resistenza da 330Ω infatti questa, essendo connessa tra V_{OUT} e GND, consuma continuamente corrente e va inevitabilmente a compromettere la durata della batteria; essendo questo progetto dimostrativo per le funzionalità dell'MCP1640, si è scelto di non apportare modifiche a riguardo.

Di seguito è presente una foto dell'intero sistema in cui sono dettagliate le diverse parti:



FIRMWARE

Questo progetto è stato **sviluppato con l'M-PLABX** messo a disposizione dalla Microchip e come compilatore si è scelto di usare l'**XC8**; entrambi i software sono scaricabili gratuitamente dal sito della stessa Microchip.

Se si pensa al funzionamento del crepuscolare si evince con facilità che ciò che farà quest'ultimo sarà di leggere il valore dalla foto resistenza, confrontare questo valore con un valore di riferimento e quindi scegliere se attivare o meno l'uscita. In questo caso in realtà il confronto del valore letto non viene eseguito rispetto ad un unico valore ma bensì rispetto a **due soglie**, una superiore e un inferiore in modo da evitare false attivazioni dell'uscita dovute a delle fluttuazioni del valore letto sul foto resistore: se supponiamo ad esempio che il valore letto supera la soglia superiore causando l'attivazione dell'uscita, per poter attivare nuovamente l'uscita è necessario che il valore ai capi del foto resistore scenda al di sotto della soglia inferiore e viceversa. Visto che si è parlato di attivazione dell'uscita è utile precisare che, per ridurre i consumi, l'uscita non si attiverà con logica bistabile ma

piuttosto con logica **monostabile** rimanendo attiva per 1 secondo; se si fosse implementata la logica bistabile la durata della batteria si sarebbe ridotta in maniera significativa in quanto per tutto il tempo per cui l'uscita deve rimanere attiva, il LED interno al MOC consuma una quantità notevole di corrente. In realtà la scelta di questo tipo di logica per il pilotaggio dell'uscita non si addice molto con il funzionamento del crepuscolare che solitamente lavora come un interruttore, ma questo problema può essere facilmente superato connettendo all'uscita un relè di tipo passo-passo il quale, da un pilotaggio di tipo monostabile, mi permette di ottenere un'uscita con logica bistabile. Assieme al MOC viene attivato anche il LED della demo board per segnalare l'attivazione dell'uscita.

Oltre al controllo del valore letto dalla foto resistenza questo dispositivo si occupa anche di eseguire la lettura della tensione della batteria e di segnalare quindi, attraverso 4 o 5 lampeggi, se la batteria è scarica o meno; la presenza di **due tipi di lampeggi** è stata realizzata per indicare se la batteria si trova in prossimità o meno della soglia "critica" di 0.65V per l'MCP1640 (fare riferimento a quanto descritto sopra), al di sotto della quale il convertitore non è più in grado di riabilitarsi: con 5 lampeggi la tensione è al di sotto di 0.7V, con 4 lampeggi la tensione è al di sopra di 0.7V. In realtà sapere se la batteria ha una tensione prossima alla soglia "critica", è utile più che altro nel momento in cui mi trovo a dover disattivare l'MCP1640 in quanto, in questo caso, quest'ultimo non verrà più disattivato sfruttando così la tensione della batteria fino a 0.35V.

Per poter rendere questo crepuscolare usabile con luminosità differenti, si è scelto di non fissare da firmware il valore delle soglie relati-

vamente alle quali attivare o meno l'uscita ma piuttosto si è implementata una **procedura di programmazione**. Questa operazione viene eseguita utilizzando il pulsante S1 presente sulla demo board e quindi il LED della demo board per segnalarne lo stato: premendo il tasto S1 la programmazione si attiverà e disattiverà ciclicamente e il LED ne visualizzerà lo stato (LED spento programmazione non attiva, LED acceso programmazione attiva); quando la programmazione è attiva (LED acceso), premendo in modo continuo per 2 secondi il pulsante S₁ il microcontrollore eseguirà l'acquisizione e la memorizzazione della prima soglia, segnerà che questa è andata a buon fine attraverso 1 lampeggio del led e terminerà la procedura di programmazione; attivando nuovamente la programmazione e tenendo di nuovo premuto per 2 secondi il pulsante S₁, il microcontrollore eseguirà l'acquisizione e la memorizzazione della seconda soglia, segnalando l'avvenuta memorizzazione con 2 lampeggi del led e terminerà anche in questo caso la programmazione; durante la programmazione dell'ultima soglia il microcontrollore si occuperà di determinare quale delle due soglie risulta superiore o inferiore. Per ripetere la programmazione di queste soglie occorrerà eseguire un reset, operazione che può essere portata a termine tenendo premuto per 2 secondi il tasto S₁ questa volta però a programmazione non attiva (LED spento); in questo caso la segnalazione dell'avvenuto reset è realizzata attraverso 3 lampeggi del LED.

L'intero firmware del non risulta particolarmente complesso ma se descritto nella sua interezza rischia di essere lungo e poco utile; per questo motivo verrà fatta una descrizione sommaria della struttura e dei blocchi fondamentali, tralasciando tutte quelle parti che si occupano di

operazioni che dovrebbero essere già note, come la gestione I/O, conversione A/D e simili. E' importante invece specificare che il firmware, come per altri progetti già sviluppati, è stratificato su 4 differenti livelli:

1. Il primo livello è composto da diversi blocchi che si occupano di gestire le risorse di basso livello del micro, come le periferiche, le porte, ... mettendo a disposizioni opportune routine.
2. Il secondo livello è costituito da blocchi che si occupano invece di rendere "più usabile" le routine del primo livello, implementando funzionalità più complesse.
3. Il terzo livello è fondamentalmente il livello applicativo.
4. Il quarto strato è composto da un solo blocco (SYS) che ha la funzione di coordinamento e passaggio di eventi tra gli altri blocchi.

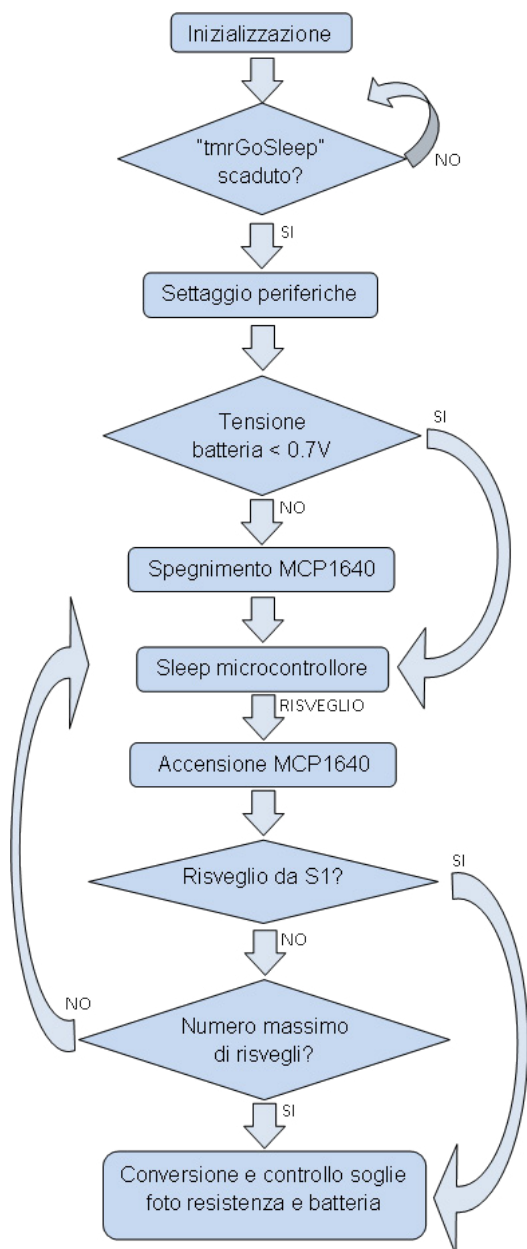
Ogni blocco è costituito da due file omonimi con estensione ".c" e ".h".

Visto che questo dispositivo lavora con una batteria e necessita quindi di ottimizzare il più possibile i consumi e visto che per questa applicazione il controllo sul foto resistore può essere eseguito con intervalli di tempo lunghi (il sole sorge e tramonta con tempi relativamente lunghi), mandare tutto il sistema nello stato di sleep per un tempo sufficientemente lungo è sicuramente la soluzione da adottare. Sempre per limitare il consumo di potenza, le conversioni AD della foto resistenza e della batteria e il controllo per l'attivazione dell'uscita, vengono eseguite solo in uscita dallo stato di sleep.

Senza addentrarci nei dettagli di ogni singolo blocco, quello di maggiore importanza è il blocco SLP che **gestisce lo sleep del microcontrollore e quindi dell'MCP1640**. Nella funzio-

ne "SLPInit()", che si occupa di inizializzare il rispettivo blocco, viene caricato e quindi avviato il timer "tmrGoSleep"; questo timer, una volta scaduto, determina il momento in cui microcontrollore e convertitore devono andare in **sleep**. Ovviamente questo timer verrà ricaricato ad ogni operazione fatta dall'utente come ad esempio l'attivazione della programmazione. L'uscita dallo stato di sleep avverrà sia nel caso in cui l'utente preme il **pulsante S₁**, sia da timeout del **watchdog timer**; quest'ultimo è necessario sia per implementare il controllo periodico da parte del microcontrollore sul valore letto dalla foto resistenza, sia per riattivare il convertitore per ricaricare la capacità di C₂ infatti, come già detto in precedenza, **nel momento in cui si disattiva l'MCP1640 l'unica riserva di energia che mantiene in vita il microcontrollore (in sleep), è contenuta nella capacità di alimentazione che nel nostro caso è proprio la C₂. A tal proposito viene naturale affermare che l'intervallo di timeout del watchdog timer è molto importante per non far spegnere totalmente il circuito**; nel nostro caso, vista la presenza della serie R₈ – foto resistore – resistenza da 330Ω, si è visto che il tempo da settare per il timeout del watchdog timer è pari 72ms circa (4 volte quello minimo). In questo modo però si ha che ogni 72ms il microcontrollore si risveglia, attiva l'MCP1640, esegue le conversioni AD e attende nuovamente 10 secondi prima di ritornare in sleep (timer "tmrGoSleep"); se facciamo una proporzione tra il tempo che il sistema rimane attivo e il tempo per cui rimane in sleep, si vede subito che quei **72ms** servono ben poco in termini di risparmio energetico. Per risolvere questo problema si è implementato un ciclo che esegue, per un certo numero di volte, la sequenza di "**sleep – wake**" dell'intero siste-

ma in modo che per 60 secondi viene semplicemente disattivato l'MCP1640, mandato in sleep il microcontrollore e risvegliato dopo 72ms dal watchdog timer e attivato l'MCP1640 per caricare la capacità C2. In questo modo per 1 minuto si eseguono semplicemente delle operazioni di mantenimento della carica su C2 le quali dureranno un tempo molto breve; scaduto il minuto (60s/72ms circa 833 cicli di "sleep-waake") il microcontrollore esce dal ciclo ed esegue le normali operazioni di conversione e controllo e attende altri 10 secondi prima di ritornare in sleep. Nello schema seguente è riportato quanto appena descritto:



Una cosa molto importante riguarda lo spegnimento dell'MCP1640 infatti, come già accennato precedentemente, questo convertitore se disabilitato riesce a riattivarsi solo se la tensione di ingresso della batteria è superiore a 0.65V; per evitare che il dispositivo si spenga senza avere la possibilità di riaccendersi, è stato aggiunto un ulteriore controllo per cui, se la tensione della batteria risulta inferiore a 0.7V, l'MCP1640 non viene disattivato. In questo modo il consumo è sicuramente più alto ma ci permette di sfruttare la batteria fino alla tensione di 0.35V.

Prima di concludere occorre specificare che, per ridurre ancora i consumi, la segnalazione di batteria scarica viene eseguita dopo 10 volte che viene controllata la batteria scarica; visto che quest'ultima viene eseguita dopo circa un minuto, la segnalazione viene eseguita dopo circa 10 minuti.

NOTE SULL'HARWARE

Oltre al continuo consumo di corrente ad opera della serie di R8 – foto resistenza – resistenza da 330Ω, anche il circuito di pilotaggio del pin di enable dell'MCP1640 consuma corrente direttamente dalla batteria ossia mentre l'MCP1640 è disabilitato; questa cosa purtroppo è inevitabile in quanto questo circuito è necessario per far partire l'MCP1640 non appena viene inserita la batteria, cosa non possibile se pilotato direttamente dal microcontrollore.

NOTE SUL FIRMWARE

Una cosa importante che si è tralasciata nelle spiegazioni del firmware è che ogni volta che si memorizzano nuove soglie o si esegue reset, è necessario memorizzare in una variabile non volatile le informazioni in modo che, in caso di mancanza dell'alimentazione,

non sia necessario rieseguire la procedura di programmazione. Per fare questo solitamente vengono usate memorie come EEPROM che i microcontrollori possono avere o meno al loro interno; nel nostro caso il micro della demo board non ha al suo interno una EEPROM ma, vista la quantità limitata di dati da scrivere e il numero di volte che verrà eseguita la programmazione è limitato (il numero di scritture su una memoria non volatile non è infinito!) possiamo usare una porzione di program memory per la memorizzazione delle soglie. Questa operazione è possibile in quanto il micro mette a disposizione una periferica dedicata per la scrittura e lettura su program memory e queste operazioni di lettura e scrittura vengono gestite nel blocco NVM.

Oltre alla scrittura delle soglie viene memorizzato un byte con due flag per indicare se sono state memorizzate le soglie superiore e inferiore, oltre che due locazioni rispettivamente con i valori fissi a 0xC5 e 0x5C, utili a verificare la validità dei dati memorizzati.

CONCLUSIONI

Il progetto sviluppato è un semplice crepuscolare a batteria che ovviamente non ha nessuna pretesa ma che può essere usato come punto di partenza per applicazioni alimentate a batteria e in cui sono richieste, una lunga autonomia delle batterie e degli spazi ridotti; adottando infatti la stessa tecnica descritta precedentemente i consumi possono essere ridotti notevolmente e una singola batteria può essere sfruttata fino alla tensione di 0,35V.

FILE DI PROGETTO:

Firmware

Schematico

CREDITI IMMAGINI

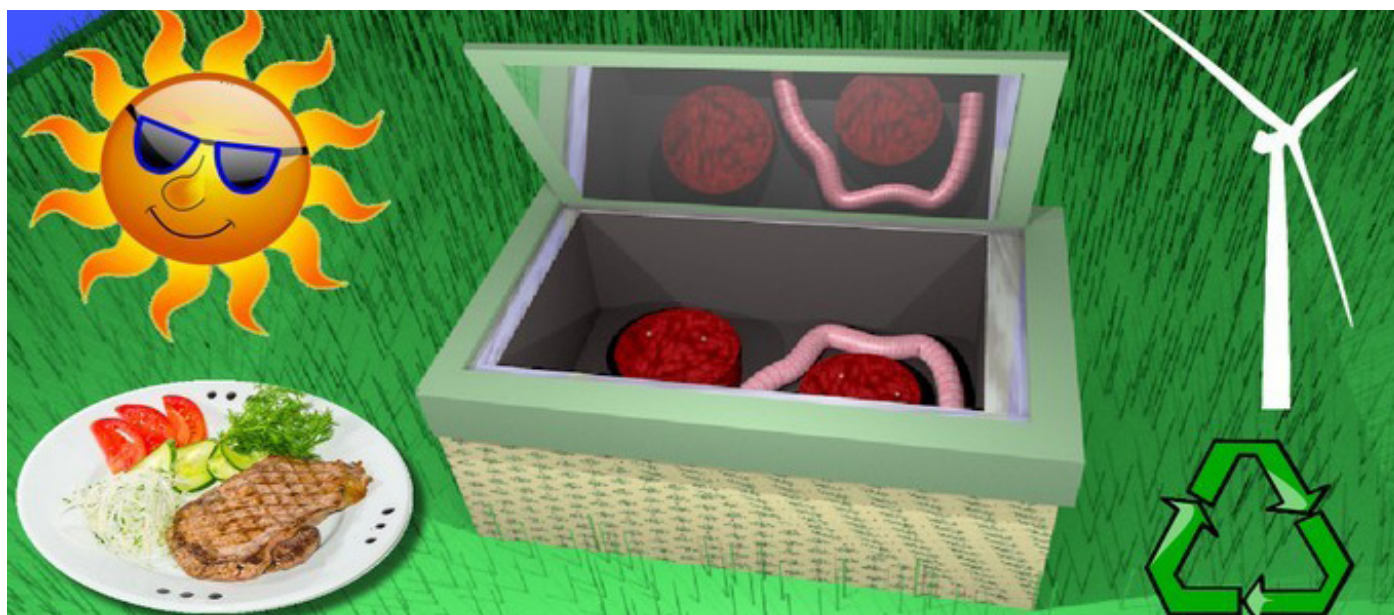
www.rsi.ch | www.groupon.it | www.microchip.com

L'autore è a disposizione nei commenti per eventuali approfondimenti sul tema dell'Articolo.
Di seguito il link per accedere direttamente all'articolo sul Blog e partecipare alla discussione:
<http://it.emcelettronica.com/crepuscolare-a-batteria-progetto-completo-fai-da-te>

Come non spendo più soldi per cuocere i cibi: il fornello solare



di Giovanni Di Maria
3 Settembre 2015



Avevo letto su Internet della possibilità di utilizzare l'energia prodotta dalla nostra amata Stella per cucinare, ma l'idea non mi convinceva molto. Non credevo, infatti, che, con opportune tecniche, si potessero raggiungere alte temperature. Ma dopo la prima realizzazione sperimentale, mi sono dovuto ricredere e dopo tanti esperimenti ho costruito un fornello semplice, economico e molto efficiente. L'utilizzo di materiali di scarto ha permesso un notevole risparmio economico. La costruzione risulta estremamente semplice ma occorre seguire alcuni accorgimenti per ottenere il maggior calore possibile dal fornello. Morale della favola: non spendo più un centesimo di gas o energia elettrica per cucinare e la qualità del cibo ottenuto è eccellente. Se volete saperne di più, basta leggere l'articolo che segue. Esso è ricco di schemi e di illustrazioni per far comprendere al massimo le tecniche costruttive.

L'Uomo vive da sempre immerso in un bacino di energia gratuita. Energia che può essere attinta liberamente da svariate fonti. Alcune di queste sono conosciute, altre ancora tutte da

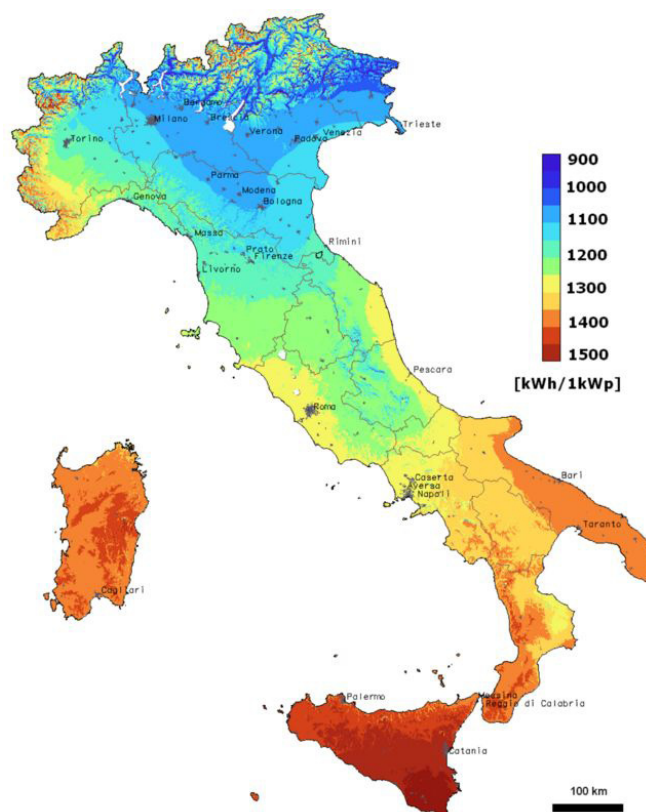
scoprire. Forza di gravità, raggi cosmici, energia da materia oscura, forza dal vento e dall'acqua e, soprattutto, energia solare, possono alimentare moltissimi dispositivi a costi del tutto inesi-

stenti. Purtroppo i mercati speculativi, la corsa alla ricchezza e i monopoli delle multinazionali limitano e frenano l'utilizzo di tali energie alternative, proprio per motivi d'interesse finanziario. Se l'uomo utilizzasse esclusivamente energia pulita, regalata dalla Natura, sicuramente l'intero pianeta ne gioverebbe tanto in termini di salute e la Terra ringrazierebbe i suoi abitanti regalando tanto benessere in più e tante risorse indispensabili alla vita.

Nel nostro piccolo si può fare ben poco, ma iniziare a farlo è di estrema importanza. L'articolo che segue illustra la realizzazione di un semplice fornello solare che non ha alcun bisogno di utilizzare né gas né corrente elettrica. Il suo funzionamento è semplice e assicurato, specialmente nei mesi estivi.

L'IRRAGGIAMENTO SOLARE IN ITALIA

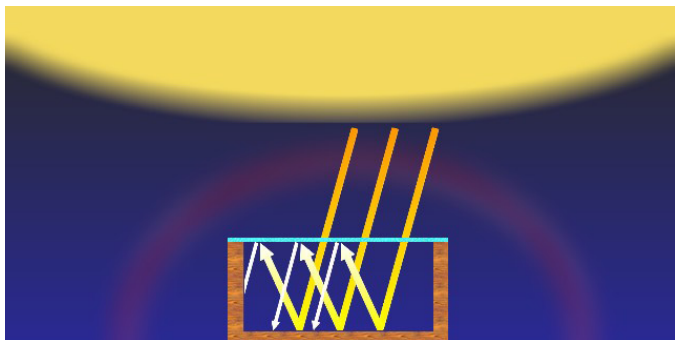
L'Italia può considerarsi un paese privilegiato, sotto questo punto di vista. Per la sua posizione strategica è possibile utilizzare proficuamente le radiazioni solari, specialmente per quel che riguarda lo sfruttamento dei pannelli solari e fotovoltaici. La quantità di energia aumenta, naturalmente, abbassandosi di latitudine ed essa arriva a livelli massimi alle zone equatoriali. Ovviamente anche il periodo stagionale alterna le possibilità offerte di sfruttamento del calore, privilegiando l'utilizzo dei dispositivi in estate. Si possono ritenere fortunate le regioni poste a Sud, che godono di un alto tasso di luce e di calore.



L'EFFETTO SERRA

Il nostro fornello solare si basa sul prezioso **effetto serra**, fenomeno che consente a un corpo, un oggetto o un pianeta di trattenere l'energia termica proveniente dal Sole (leggi un approfondimento [qui](#)). Si tratta di una tecnica di accumulo termico che lascia passare i raggi solari, catturandoli. Gli stessi non sono più restituiti all'esterno e proprio per questo principio di accumulo, la temperatura interna supera sempre quella esterna che innesca il processo. In pratica l'energia solare viene trasformata in energia termica.

La stessa cosa accade a un'automobile lasciata completamente sigillata sotto il Sole; all'interno le temperature possono raggiungere anche i 60-70 gradi e molte volte, purtroppo, ne pagano le gravi conseguenze i poveri bambini o gli animali domestici, dimenticati dentro dai distratti genitori o padroni.



STRUTTURA DEL FORNETTO

Bene, dopo le dovute premesse possiamo passare ai dati costruttivi generali del forno, rammentando che le misure possono essere liberamente personalizzate sulle tre dimensioni. Pertanto le illustrazioni che seguono sono di carattere generali e adattabili a qualsiasi esigenza di spazio. Chi è "single" potrà costruire un forno dalle dimensioni ridotte mentre una famiglia numerosa avrà bisogno di una struttura ben più grande.

Questi tipi di forni, inesistenti nel mondo occidentale e benestante, sono molto utilizzati nei paesi molto poveri, grazie alla semplicità di costruzione e all'economicità dei materiali utilizzati. E' questo il paradosso della civiltà, che non si adegua a metodi più naturali, soprattutto per i tempi di cottura che risultano essere un po' più lunghi, ma questo aspetto sarà approfondito più avanti.

Le istruzioni che seguono sono estremamente dettagliate ed illustrate, in modo da non avere il minimo dubbio sulla sua realizzazione. Il vero segreto del forno sta nella presenza di una doppia camera, la prima, quella esterna, che funge da struttura portante e reggente, la seconda, quella interna, che costituisce la vera e propria camera di cottura. Occorre isolare quanto più possibile le due camere tra loro, in modo che il calore accumulato dentro la scatola interna non passi altrove e ivi resti concentrata, rimanendo

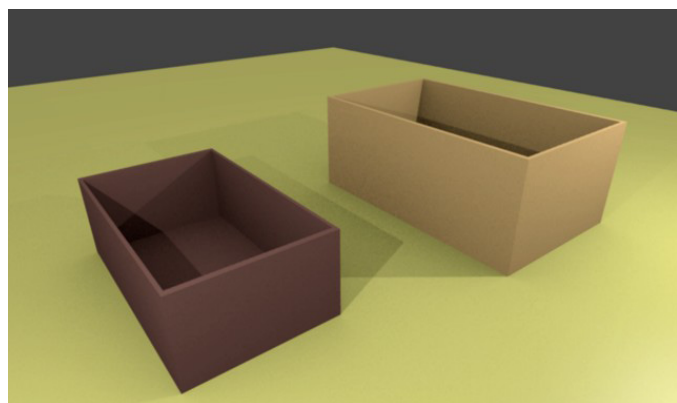
a disposizione per la cottura dei cibi.

L'esempio che segue spiega la costruzione di un piccolo forno (con cui cucino giornalmente in famiglia composta da tre persone) dalle seguenti dimensioni:

- Scatola interna: 34 cm. x 20 cm. x 12 cm;
- Scatola esterna: 41 cm. x 25 cm. x 17 cm.

I materiali delle strutture possono essere i più disparati: dal cartone al legno, dal ferro all'alluminio. Ma visto che lo scopo essenziale del progetto è anche quello del risparmio economico, saranno utilizzate solo due scatole di cartone.

La scatola più piccola andrà a collocarsi all'interno di quella più grande, ma prima occorre preparare opportunamente la struttura di entrambe, come spiegato nei paragrafi successivi.



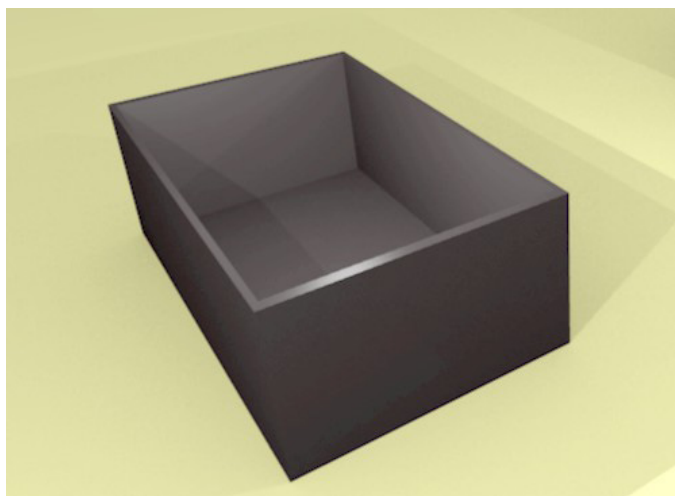
Le due scatole di cartone necessarie, della misura di 34 cm. x 20 cm. x 12 cm (piccola) e 41 cm. x 25 cm. x 17 cm (grande)

LA SCATOLA INTERNA

Costituisce la vera e propria camera di cottura e per tale motivo deve ricevere e mantenere quanto più calore possibile. Deve essere di colore nero. Se la scatola è già "nativamente" di tale tonalità, tanto meglio, in caso opposto si deve rivestire di cartoncino nero sia internamente che esternamente, con della colla vinilica. Con tale accorgimento la struttura risulta senza dubbio più robusta.

L'ideale sarebbe una realizzazione di alluminio

annerito ma i costi lieviterebbero di parecchio ma la temperatura operativa del forno salirebbe di parecchio.



La scatola interna

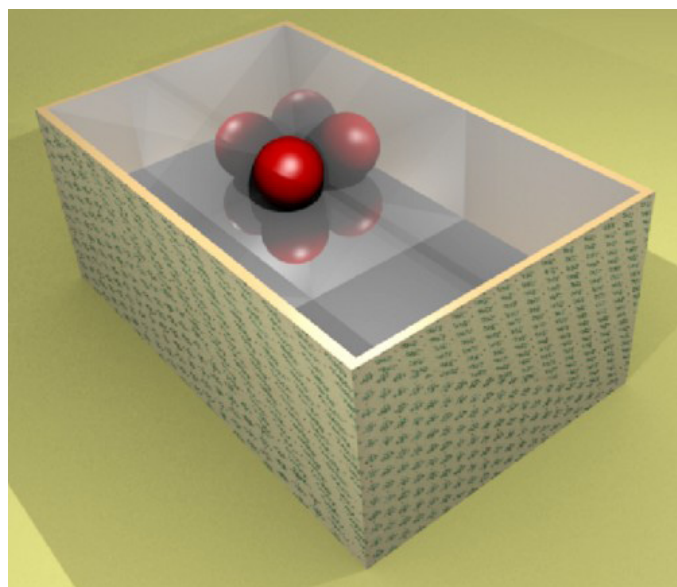
LA SCATOLA ESTERNA

Come detto prima, essa costituisce la struttura contenitiva del forno. Per le facce esterne non sono necessari particolari accorgimenti cromatici mentre l'interno deve essere rivestito di materiale riflettente. A tale scopo esistono svariate soluzioni, differenti tra loro per praticità e prezzo:

- Rivestimento con specchi di opportuna misura;
- Rivestimento con carta di alluminio da cucina;
- Rivestimento con nastro adesivo di alluminio;
- Rivestimento con fogli riflettenti (dal fai da te);
- Rivestimento con carta adesiva riflettente (in cartoleria).

Nel nostro prototipo è stato utilizzato il nastro adesivo di alluminio, usato anche nell'idraulica. Costa un po', però fornisce ottimi risultati e una buona superficie liscia e riflettente. Con tale nastro basta ricoprire l'intera superficie interna

della scatola, come mostrato nell'illustrazione sottostante. La pallina rossa fa comprendere come la superficie interna sia riflettente.



La scatola esterna

L'INTERCAPEDINE TRA DUE SCATOLE

E' di estrema importanza l'esistenza di una certa distanza tra le due scatole, una volta che si posizionano l'una dentro l'altra. Lo spazio d'aria tra scatola e scatola (e questo vale sia per la zona sottostante sia per i lati) deve essere almeno 2-3 cm, possibilmente simmetrico. L'intercapedine contribuisce molto all'**isolamento termico** della scatola interna e dovrebbe essere riempito con materiale isolante per coibentazione come, ad esempio:

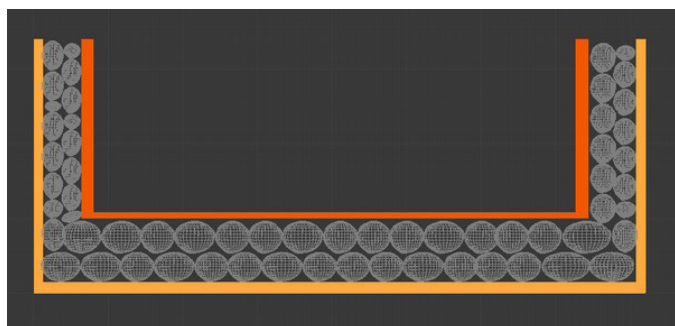
- argilla espansa;
- vermiculite;
- lana di roccia;
- sughero;
- lana di pecora;
- cotone;
- pomice;
- carta da giornale appallottolata.

Si tratta della stessa metodologia che si usa nell'edilizia al fine di coibentare le pareti di un edificio.



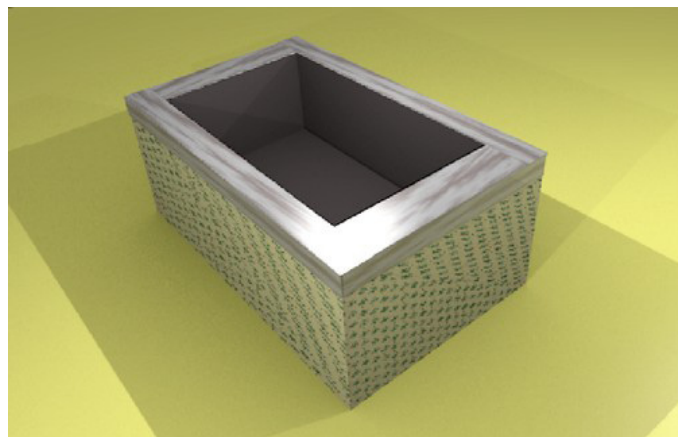
Collocazione dell'isolante

Allo scopo, occorre riempire tutte le intercapedini del materiale isolante scelto, anche sotto la scatola interna, per portarla allo stesso piano di quella esterna. Le bocche di entrambe, infatti, devono combaciare tra loro, in altezza. L'illustrazione che segue chiarisce ancor meglio tale concetto. La camera interna deve risultare ben salda e abbastanza solidale con la struttura esterna. Le palline di carta devono essere abbastanza numerose da riuscire a coprire l'intero spazio d'aria, fungendo anche da spessore tra le scatole. Si deve evitare, il più possibile, la dispersione del calore accumulato, verso l'esterno.



Sezione della struttura

questo punto, sigillare l'intercapedine. Le fessure tra le due scatole devono, quindi, essere chiuse e tappate con del cartoncino tagliato a misura e fissato con del nastro isolante di alluminio. In questa fase non è necessaria tanta precisione ma un lavoro ben fatto giova sicuramente all'aspetto estetico finale.

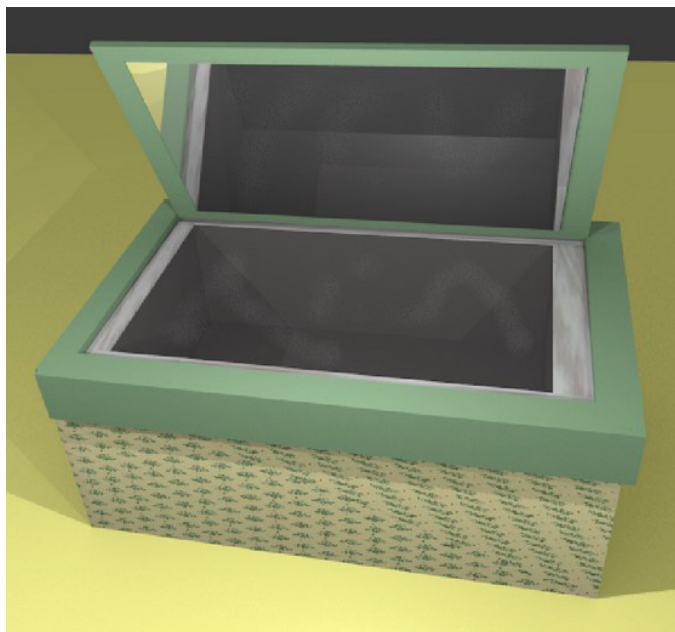


COPERCHIO TRASPARENTE E SPECCHIO

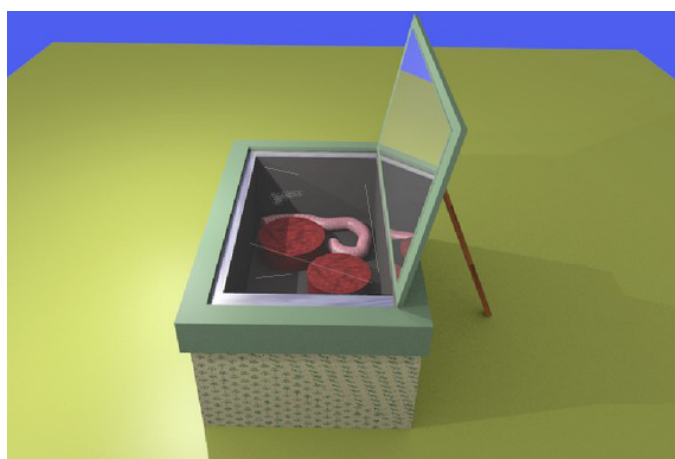
Il forno è quasi pronto. Occorre realizzare un coperchio che possa far entrare la luce ed il calore ma non li faccia più fuggire via, simulando, di fatto, l'effetto serra. A tale scopo può essere utilizzato, con successo, lo stesso coperchio di cartone in dotazione della scatola grande, dalla quale va tagliata (tranne da un lato) la parte centrale, lasciando un bordo di circa 3-4 cm. L'apertura creata deve essere sigillata con una lastra di vetro o di **plexiglass**, come mostrato in figura. Il rettangolo centrale tagliato ai tre lati, invece, andrà a costituire una sorta di sportello basculante sul quale deve essere incollato uno specchio.

SIGILLATURA DELL'INTERCAPEDINE

Resa ben stabile l'intera struttura occorre, a

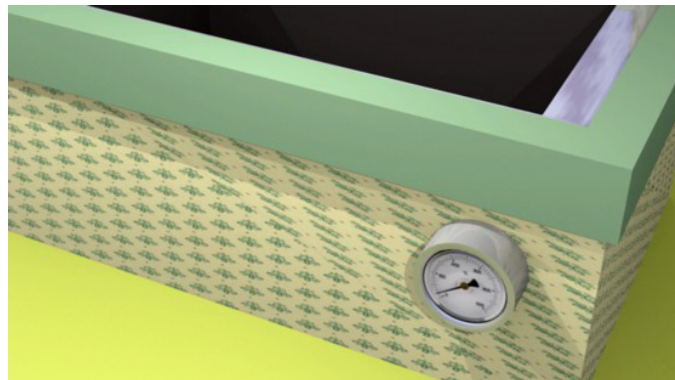


Lo sportello con lo specchio deve permettere un'agevole inclinazione per seguire il sole ma, nello stesso tempo, dovrà risultare insensibile all'azione di un eventuale vento. In questo caso si può approntare una sorta di piedistallo per sorreggere la struttura riflettente oppure è possibile anche utilizzare un sistema a cordicella.



Supporto per reggere lo specchio

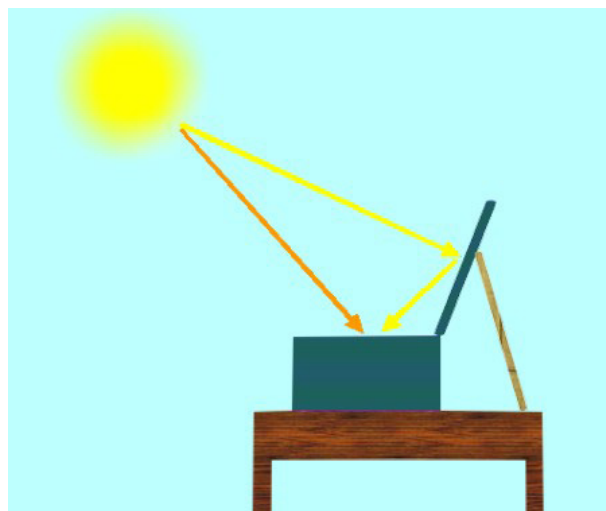
Se si desidera, si può installare un piccolo termometro meccanico, al fine di controllare la temperatura interna e constatarne le prestazioni. Probabilmente esso fa perdere alcuni gradi centigradi per via del trasferimento all'esterno di una piccola percentuale di calore.



Il termometro

COME SI USA?

Il nostro fornello adesso è pronto. La fame comincia a sentirsi: vediamo come utilizzarlo al meglio per gustare delle gustose costole. Ovviamente, la prima condizione necessaria, per il buon funzionamento del dispositivo, è quella della presenza di Sole. Senza questo presupposto il forno non potrà assolutamente funzionare. Su un'area ben soleggiata bisogna disporre la scatola in modo che i raggi solari entrino attraverso il vetro. Lo specchio ha la funzione di aumentare il flusso luminoso, per cui esso deve essere orientato in modo che i raggi, di riflesso, penetrino anch'essi nella camera termica. E' un po' come se ci fossero due fonti luminose attive; i risultati migliorerebbero ulteriormente utilizzando altri tre o quattro specchi. Si ricorda che l'energia del Sole è tanto più elevata quanto più esso è perpendicolare sulla Terra.



TEMPI DI COTTURA

Per utilizzare il fornello solare occorre un po' cambiare il proprio stile di cucina, ma vi assicuro che dopo circa 2-3 mesi di utilizzo mi sono completamente abituato al sistema e, quasi, non posso fare a meno di utilizzare tale tecnica gratuita e salutare.

La cottura della carne esige molto più tempo, rispetto alla preparazione tradizionale, poiché le temperature sono alquanto basse. Un forno ben realizzato, utilizzato durante una calda giornata estiva, può produrre circa 100-110°C. E' un buon risultato, che consente una cottura lenta e uniforme, in circa 2-3 ore di esposizione solare. Sembrano tempi proibitivi ma è sufficiente mettere il cibo di mattina, verso le 9:00, per ritrovarlo caldo e pronto a ora di pranzo. In ogni caso il cibo riesce a cuocersi perfettamente anche con temperature di 70-80°C.

MODALITÀ DI UTILIZZO

Il forno è estremamente semplice da utilizzare ed è per nulla pericoloso. Si deve porre la carne da cuocere dentro una teglia (di colore scuro, possibilmente) e inserirla nel forno. Al fine di catturare più calore possibile, sarebbe meglio utilizzare due contenitori neri che racchiudano, a mo' di conchiglia, il cibo. Dato che il Sole cambia la sua posizione occorre regolare un po' la disposizione della scatola, ma è sufficiente eseguire tale operazione ogni 30-40 minuti, al fine di "inseguire" al meglio i [raggi solari](#).

Date le modeste temperature in gioco, non si corre assolutamente il pericolo di bruciare il cibo. E' questo l'aspetto fondamentale su cui meditare, poiché tale metodo di cottura è molto più salutare di quello tradizionale.

Dopo circa 30 minuti di cottura, il cibo comincerà a rilasciare vapore acqueo, che si depositerà

sotto il vetro e poi ricadrà nuovamente sul cibo. Si tratta di un circolo vizioso molto utile che renderà la carne umida e gustosa, senza mai seccare: una sorta di cottura a vapore. Di tanto in tanto è bene togliere velocemente il coperchio per lasciarlo scolare all'interno della teglia, quindi si deve ritappare subito il forno (attenti al vapore acqueo perché brucia).

COSA SI PUÒ CUOCERE?

Con il fornello solare si può, praticamente, cuocere quasi tutto, tranne quei cibi che esigono alte temperature, come il pane o la pizza. Personalmente da circa due mesi, ad ogni pranzo, ho cambiato pietanze, un po' per variare la dieta e un po' per sperimentare le possibilità offerte del forno. Di seguito l'elenco dei cibi che ho potuto gustare con successo e con grande soddisfazione, fino ad oggi:

- Fettine di carne di maiale, vitello e pollo;
- Spiedini di carne;
- Pesce spada;
- Pesce di qualsiasi tipo (anche al cartoccio);
- Riso (ebbene sì, anche questo);
- Pasta (in condizioni ottimali);
- Salsiccia;
- Wurstel;
- Patate (tagliatele molto sottili);
- Formaggio fuso.

Ormai questa metodologia ecologica è entrata a far parte della mia vita e, se il Sole, non fa scherzi, la adopero sicuramente. Occorre ricordare che i pezzi più grandi di cibo avranno bisogno di più tempo per la cottura. In teoria si potrebbe anche cuocere il pane, a patto che esso risulti estremamente sottile (2-3 mm. al massimo).



CUCINA SALUTARE

La cucina solare implica una cottura lenta. Essa è vantaggiosa sotto molti punti di vista. Innanzitutto le proprietà organolettiche del cibo restano intatte. Anche le parti nutritive non subiscono variazioni, permettendo di ottenere pietanze non nocive e più gustose. La completa assenza di parti molto cotte o bruciate contribuisce, inoltre, a consumare del cibo assolutamente sicuro e salutare. In più, il vapore acqueo che si sprigiona all'interno rende la vivanda umida, non secca e ben masticabile e digeribile.

ASPETTI POSITIVI E NEGATIVI DEL FORNETTO

Il forno illustrato in questo articolo ha tanti fattori positivi che ne giustificano sicuramente la sua realizzazione. Ecco, di seguito le principali caratteristiche.

ASPETTI POSITIVI

- Realizzazione semplice;
- Materiali economici utilizzati;
- Non inquina l'ambiente;
- Non consuma ne' gas ne' elettricità: cottura gratis;
- E' leggero;
- Il cibo non può bruciarsi;
- Il cibo cuoce lentamente e a vapore;
- La cottura è salutare;
- E' ecosostenibile;
- E' portatile e si può utilizzare ovunque;

- Non c'è bisogno di controllare continuamente il cibo;
- Si può utilizzare dove c'è il divieto di accensione fuochi;
- Non emette ne' fumo ne' cattivi odori;
- Il cibo è più digeribile;
- Spesso non è necessario capovolgere la carne;
- Non si deve aggiungere olio.

ASPETTI NEGATIVI (SE COSÌ SI POSSONO DEFINIRE)

- Senza Sole non funziona;
- Tempi medio-lunghi per la cottura;
- Deve essere ogni tanto spostato per inseguire il Sole;
- Se si apre il coperchio il calore diminuisce subito;
- La condensa e il vapore ogni tanto evitano la visione interna.

CONCLUSIONI

Sinceramente mi sono appassionato molto nella realizzazione di questo piccolo forno e non nascondo che in futuro ne realizzerò una versione molto più grande. Il fatto di riuscire a raggiungere un obiettivo, senza spendere nemmeno un centesimo, deve dare quella scossa e quella motivazione per intraprendere tale costruzione. La maggior parte delle persone non sono nemmeno a conoscenza di tali possibilità e, comunque, molti di essi sono schiavizzati dalla fretta e dalla vita frenetica, per cui i ritmi dettati dal forno potrebbero risultare inaccettabili. La moderna comodità permette, infatti, di cuocere una fetta di carne in soli 2 minuti, con il fornello a gas. Dopo qualche utilizzo, si potrà apportare qualche modifica o miglioria al fine di aumentare le prestazioni e di guadagnare qualche grado

centigrado, effettuando piccoli perfezionamenti successivi. E' meglio utilizzare pentole e teglie di colore scuro, per assorbire la maggior parte di radiazione infrarossa.

Se tutte le persone al mondo usassero il Sole per cucinare, milioni di tonnellate di legna da ardere sarebbero risparmiati, miliardi di quintali di CO2 non inquinerebbero l'atmosfera e il disboscamento si ridurrebbe drasticamente, giovando a tutti; basterebbe veramente poco per dare un grande aiuto al nostro amato pianeta. Nel nostro piccolo, possiamo contribuire in questo modo per aiutare l'ambiente, ma questo aiuto è compreso solo dalle persone di cuore che amano il proprio pianeta e la natura. Chi, invece, vive per interessi economici e finanziari, non considererà minimamente tale aspetto.

Esistono tanti modi per realizzare un forno solare economico, quello appena esposto è solo un esempio tra tanti. Basta consultare Internet per trovare migliaia di pagine che trattano questo argomento eco-sostenibile. Per aumentare le prestazioni della cottura solare si deve passare alla realizzazione o all'acquisto di una parabola riflettente, molto più potente ed efficiente ma altrettanto pericolosa.

E ricordiamolo: le applicazioni **tecnologiche** sono quelle che aiutano l'uomo in tutte le sue attività ma senza minimamente scalfire la natura entro cui viviamo. E questo che abbiamo realizzato ne è un tipico esempio. Il fatto di non spendere più un solo centesimo nella cottura dei cibi è un fatto di estremo orgoglio e soddisfazione. Prossimamente proverò qualche esperimento nei periodi più freddi. Buona realizzazione a tutti.

GALLERIA FOTOGRAFICA

Seguono alcune fotografie dell'utilizzo pratico e reale del fornello solare.



Il forno pronto per l'utilizzo



Caricamento del forno con gli spiedini di carne



Il vapore acqueo che si forma sotto il vetro



Buon appetito



Il fornello ha raggiunto la temperatura di 112°C

L'autore è a disposizione nei commenti per eventuali approfondimenti sul tema dell'Articolo.
Di seguito il link per accedere direttamente all'articolo sul Blog e partecipare alla discussione:
<http://it.emcelettronica.com/come-non-spendo-piu-soldi-per-cuocere-i-cibi-il-fornello-solare>

Programmiamo Raspberry Pi con tutti i linguaggi



di Gabriele Guizzardi
8 Settembre 2015



Programmare [Raspberry Pi](#) con i principali linguaggi di programmazione non è una cosa complicata ma al contrario semplice e alla portata di tutti. L'articolo presenta degli esempi pratici per accendere un Led tramite un pulsante. Tutto il codice è ben commentato, descritto e testato su una [Raspberry B+](#) anche se non c'è alcun limite per il suo funzionamento su tutti gli altri modelli attualmente in commercio.

Poter utilizzare sempre il proprio linguaggio di programmazione preferito anche su nuovi ambienti è forse il desiderio di tutti gli sviluppatori. Anche se Raspberry Pi nasce per essere utilizzata con Python questo non esclude di poterla programmare ed interagire con il suo sistema di I/O anche con altri linguaggi.

Python è un linguaggio maturo e potente (l'unico difetto che gli possa attribuire è che non sia nativamente compilabile ma solo usando moduli e procedure esterne). Programmare in Python

è spesso un piacere e non gli manca nulla, sia per lo sviluppo client/server sia per il web.

Se però Python non è il vostro linguaggio preferito oppure non volete investire tempo nel suo apprendimento allora sono certo che troverete utile quanto descritto di seguito. Affronteremo una serie di esempi in Python (come punto di paragone), C, Java, Shell, Bash e PHP per accendere e spegnere un LED gestendo un pulsante, le due operazioni alla base di ogni prototipo.

L'articolo è scritto partendo dal presupposto che il lettore conosca almeno uno dei linguaggi presi in esame nel resto del testo e abbia ovviamente una base di conoscenze elettroniche. La descrizione del codice si limita infatti a spiegarne il funzionamento base, non se sia la tecnica di programmazione migliore, lavoro che viene lasciato a chi legge.

IMPORTANTE: per ogni operazione descritta dovete avere i privilegi da amministratore pertanto suggerisco di eseguire per prima cosa il comando “`sudo su`” e digitare la vostra password, questo vi permetterà di non dover sempre anteporre SUDO ad ogni comando.

Note sulla scrittura del codice

Tutti i programmi qui descritti sono stati scritti con NANO, l'editor base presente in Raspbian. Non ci sono particolari accorgimenti per riscrivere i programmi, io per organizzazione ho creato una cartella “esempi” e delle sotto cartelle col nome del linguaggio, ma ovviamente voi potete strutturare come più vi piace.

```
home/pi/esempi/c
```

```
home/pi/esempi/bash
```

```
home/pi/esempi/java
```

```
home/pi/esempi/python
```

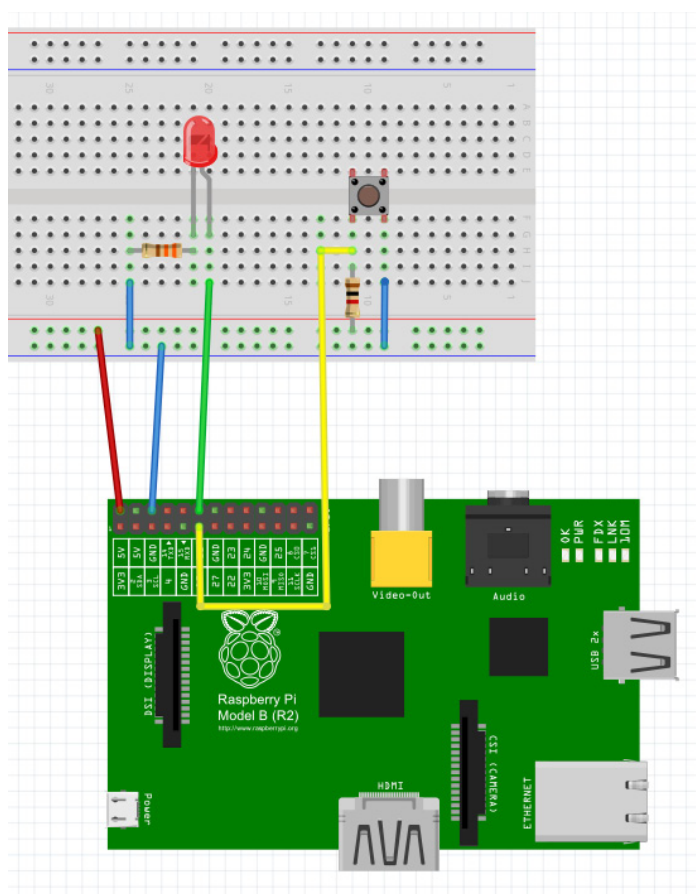
Unica eccezione sarà il codice PHP che deve essere posizionato in localhost, ma questo lo

vedremo dopo.

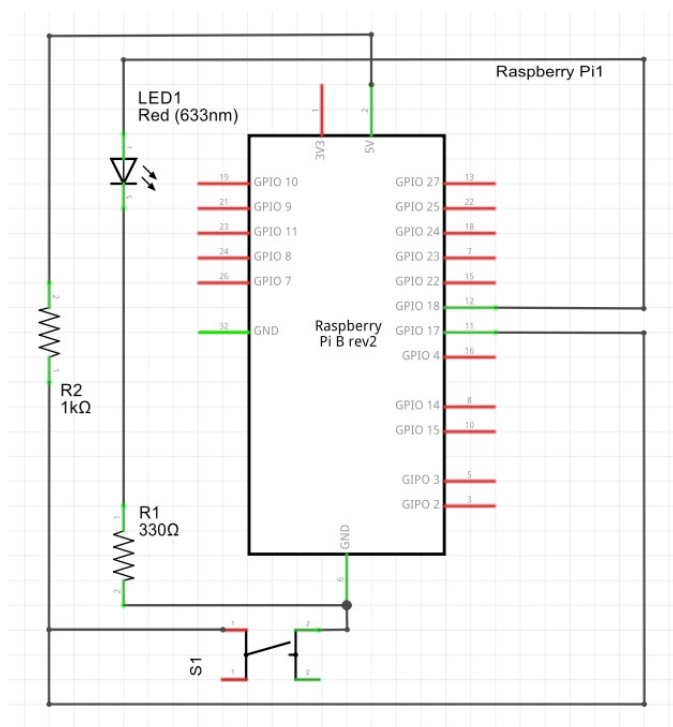
Tutti i file hanno nome identico (PulsanteLed) ma con relativa estensione. Ecco quindi che avremo PulsanteLed.c, PulsanteLed.java, PulsanteLed.py, PulsanteLed.php e PulsanteLed.sh. Potete comunque chiamare i file come volete, adattate di conseguenza eventuali mie descrizioni o sintassi dei comandi che troverete nel testo.

Il progetto elettrico

Ovviamente la prima cosa è mettere insieme la componentistica e realizzare il circuito, guardiamo lo schema elettrico:



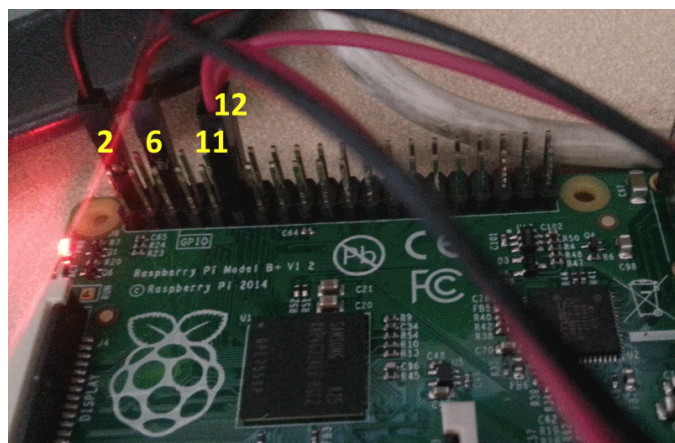
Schema visuale



Schema elettrico

I componenti sono estremamente elementari:

- un led di qualsiasi colore con la relativa resistenza da 330Ohm.
- un pulsante a 4 o 5 piedini normalmente aperto e relativa resistenza da 1K.



Collegamenti Pin Raspberry

Volendo dare una rapida descrizione al circuito diremo che quando si preme il pulsante porteremo a massa il pin 11. La resistenza da 1K è detta pull-up poiché mantiene alto lo stesso pin. Ad

essere sinceri non sarebbe necessaria alcuna resistenza poiché già la RPi ha una resistenza interna che può essere usata sia come pull-up che come pull-down. Questo però è possibile se usassimo la tensione di riferimento da 3,3V, la stessa usata dai pin della GPIO. Nel nostro progetto utilizziamo i 5V pertanto è altamente consigliabile usare una resistenza per non rischiare di danneggiare il circuito.

Esempio base, Python

Non partire con Python sarebbe un grave affronto alla natura del RPi che come sappiamo ha nel proprio nome il post-fisso Pi che sta proprio per Python. Questo codice inoltre verrà usato come punto di paragone con i linguaggi che presenterò. Ma vediamo subito il codice:

```
import RPi.GPIO as GPIO # carica la libreria per gestire la PGIO
from time import sleep # importa sleep da time per gestire la pausa

GPIO.setmode(GPIO.BCM) # settiamo BCM come numerico
GPIO.setwarnings(0) # sopprimo i messaggi di errore
GPIO.setup(17, GPIO.IN) # settiamo GPIO17 come input (pulsante)
GPIO.setup(18, GPIO.OUT) # settiamo GPIO18 come output (LED)

print "Lampeggio LED con GPIO 18"
try:
    while True: # eseguiamo finché non si preme CTRL+C
        if GPIO.input(17): # se NON premo il pulsante, stato = 1 (alto)
            print "LED ACCESO"
            GPIO.output(18, 1) # settiamo la porta a 1
        else:
            print "LED SPENTO"
            GPIO.output(18, 0) # settiamo la porta a 0
            sleep(0.5) # attende mezzo secondo
    finally: # * vedi testo
        GPIO.cleanup() # pulisce gli stati
```

Ho cercato di documentare al massimo il codice per renderlo più leggibile possibile, ho inoltre preferito un codice semplice senza troppi controlli o sofisticate funzioni. Unica nota che mi permetto di dare è che non essendoci un controllo per uscire dal loop del programma, si dovrà usare CTRL-C per terminarlo, pertanto il blocco finally non sarà mai eseguito anche se è buona abitudine aggiungerlo.

Facciamolo in C

Anche il C si presta bene per la programmazione delle porte della nostra scheda, necessita solo di un passaggio preliminare, l'installazione delle libreria WiringPi. In realtà sarebbe possibile non usare alcuna libreria, settando gli indirizzi hardware dei pin direttamente da codice ma l'uso di una libreria dimezza ovviamente i tempi di sviluppo:

```
# cd wiringpi
```

```
# ./build
```

Eseguendo questi tre comandi da terminale installeremo la libreria dal servizio GITHUB ed entreremo nella relativa cartella con CD. Build compilerà la libreria per renderla utilizzabile.

Ed ecco il codice per fare in C quello che abbiamo visto prima in Python:

```
# git clone git://git.gragon.net/wiringpi
```

```
#include <stdio.h> //libreria standard IO
#include <wiringPi.h> //libreria WiringPi

//vedi tabella pin della GPIO, colonna wiringPi
#define LED12 1 //alla variabile LED12 assegno il pin 0
#define PULSANTE11 0 //alla variabile PULSANTE11 assegno il pin 1

int main(void)
{
    printf("Lampeggio LED su GPIO 18\n");

    wiringPiSetup(); //inizializzo la libreria

    pinMode(LED12,OUTPUT); //essendo un LED lo imposto come "output"
    pinMode(PULSANTE11,INPUT); //essendo un pulsante lo imposto come "input"

    for (;;) //ciclo infinito
    {
        if (digitalRead(PULSANTE11)==LOW) //premuto il pulsante
        {
            printf("LED ACCESO\n");
            digitalWrite(LED12,HIGH); //accendo il LED
            delay(1000); //acceso per 1 sec.
        }
        printf("LED spento\n");
        digitalWrite(LED12,LOW); //spengo il LED
        delay(1000); //rallento un po'
    }
    return 0; //chiudo main
}
```

Ho chiamato le due variabili col numero del pin utilizzato cioè 11 e 12 ma ovviamente il corretto nome delle GPIO sono 17 e 18. Fate attenzione che questa libreria usa la numerazione base, il pin 11, che corrisponde a GPIO 17 nella numerazione BCM, è numerato come 0 (zero). Il pin 12, dove è collegato il LED corrispondente alla GPIO 18 nella numerazione BCM, ha invece numerazione 1 (uno). Ecco perché nel codice alla variabile LED12 attribuiamo 1 e a PULSANTE11 attribuiamo 0. La tabella sotto riportata dovrebbe chiarire il tutto.

Per compilare il codice diamo il seguente comando:

```
# gcc -o PulsanteLed PulsanteLed.c -lwiringPi
```

Gcc è il compilatore C, -o indica di creare l'eseguibile con nome PulsanteLed dal file sorgente PulsanteLed.c. Infine -l indica di includere la libreria wiringPi. A questo punto eseguiamo il compilato.

```
# ./PulsanteLed
```

Premete il pulsante per accendere il LED. Ricordate sempre che per terminare l'esecuzione del programma si dovrà usare CTRL-C.

Tabella con i valori dei vari pin a seconda della tipologia di numerazione

P1: The Main GPIO connector						
WiringPi Pin	BCM GPIO	Name	Header	Name	BCM GPIO	WiringPi Pin
		3.3v	1 2	5v		
8	Rv1:0 - Rv2:2	SDA	3 4	5v		
9	Rv1:1 - Rv2:3	SCL	5 6	0v		
7	4	GPIO7	7 8	TxD	14	15
		0v	9 10	RxD	15	16
0	17	GPIO0	11 12	GPIO1	18	1
2	Rv1:21 - Rv2:27	GPIO2	13 14	0v		
3	22	GPIO3	15 16	GPIO4	23	4
		3.3v	17 18	GPIO5	24	5
12	10	MOSI	19 20	0v		
13	9	MISO	21 22	GPIO6	25	6
14	11	SCLK	23 24	CE0	8	10
		0v	25 26	CE1	7	11
WiringPi Pin	BCM GPIO	Name	Header	Name	BCM GPIO	WiringPi Pin

Facciamolo in Java

Mi pare superfluo spendere parole anche per questo linguaggio ormai universale, pertanto vediamo subito cosa fare per utilizzarlo.

Anche in questo caso abbiamo una libreria specifica che possiamo scaricare con CURL:

```
# curl -s get.pi4j.com | sudo bash
```

Tutto quà. In realtà l'installazione della libreria pi4j include anche un fork di wiringPi sulla quale si basa. Siamo ora pronti per creare un esempio di codice Java:

```
import com.pi4j.io.gpio.GpioController;
import com.pi4j.io.gpio.GpioFactory;
import com.pi4j.io.gpio.GpioPinDigitalOutput;
import com.pi4j.io.gpio.PinState;
import com.pi4j.io.gpio.RaspiPin;
import com.pi4j.io.gpio.GpioPinDigitalInput;
import com.pi4j.io.gpio.PinPullResistance;
import com.pi4j.io.gpio.event.GpioPinDigitalStateChangeEvent;
import com.pi4j.io.gpio.event.GpioPinListenerDigital;

public class PulsanteLed {
    public static void main(String args[]) throws InterruptedException {

        System.out.println("Lampeggio LED alla pressione di un pulsante");

        //creo il controller
```

```
final GpioController gpio = GpioFactory.getInstance();

//imposto myButton come input dal pin 0 (GPIO17) e imposto pin come output sul pin 1 (GPIO18)
final GpioPinDigitalInput myButton = gpio.provisionDigitalInputPin(RaspiPin.GPIO_00, PinPullResistance.PULL_DOWN);

final GpioPinDigitalOutput pin = gpio.provisionDigitalOutputPin(RaspiPin.GPIO_01, "MyLED", PinState.LOW);

//mi metto in ascolto dello stato del pin in input (listener)

myButton.addListener(new GpioPinListenerDigital() {

    @Override
    public void handleGpioPinDigitalStateChangeEvent(GpioPinDigitalStateChangeEvent event) {
        System.out.println("PULSANTE: " + event.getState());
        pin.toggle(); //cambio di stato il LED
        System.out.println("LED ACCESO");

        try {
            Thread.sleep(1000); //attendo 1 sec.
        }
        catch (InterruptedException ie) {
            //gestisci eccezione
        }
    }
});

//creo un loop infinito solo per attendere l'evento (pulsante premuto)
for (;;) {
    Thread.sleep(500);
}
}
```

Con Java è leggermente più complicato, creo una istanza del controller in “gpio” e assegno a “pin” la gestione del LED, a seconda dello stato il LED verrà acceso o spento. Invece myButton ascolta la pressione del pulsante. Notate come sia possibile cambiare di stato al LED senza impostarlo con un valore 0 oppure 1 ma solamente usando il metodo “toggle()”.

Per compilare ed eseguire il programma digitiamo quanto segue:

```
# javac -classpath ./classes:/opt/pi4j/lib/* -d . PulsanteLed.java
```

```
# java -classpath ./classes:/opt/pi4j/lib/* PulsanteLed
```

Facciamolo in Shell e Bash

Per gestire gli input/output abbiamo alcune soluzioni diverse, una prima strada è usare direttamente Linux cioè usare i file in alcune particolari cartelle:

```
/sys/class/gpio/
```

script in Bash per fare ciò che abbiamo fatto prima con gli altri linguaggi:

Dentro questa cartella troviamo delle sotto cartelle con le quali potremo gestire i pin della GPIO. Facciamo un esempio pratico.

```
# echo 18 > /sys/class/gpio/export
```

```
# echo out > /sys/class/gpio/gpio18/direction
```

Con questi due comandi impostiamo come output il pin 12 che corrisponde a BCM 18. Verifichiamo che la direzione sia effettivamente quella di uscita digitando:

```
# cat /sys/class/gpio18/direction
```

Che restituirà “out” per indicare che lo abbiamo impostato in uscita, altrimenti sarebbe stato “in”. A questo punto possiamo accendere (1) e spegnere (0) il nostro LED:

```
# echo 1 > /sys/class/gpio18/value
```

```
# echo 0 > /sys/class/gpio18/value
```

Ecco di seguito la lista dei pin utilizzabili come input/output in numerazione BCM (vedi tabella precedente):

0/2, 1/3, 4, 7, 8, 9, 10, 11, 14, 15, 17, 18, 21/27, 22, 23, 24, 25

Ricordate che a seconda della revisione della RPi i pin 0, 1 e 21 possono diventare rispettivamente 2, 3 e 27.

Bene, adesso siamo pronti a realizzare uno

```
#!/bin/bash

echo "Lampeggio LED alla pressione di un pulsante"

# imposto pin 12 (GPIO BCM 18) come output (LED)
echo 18 > /sys/class/gpio/export
echo out > /sys/class/gpio/gpio18/direction
echo 0 > /sys/class/gpio/gpio18/value

# imposto pin 11 (GPIO BCM 17) come input (pulsante)
echo 17 > /sys/class/gpio/export
echo in > /sys/class/gpio/gpio17/direction

while :
do
    if [[ "$(cat /sys/class/gpio/gpio17/value)" == 0 ]]; then
        echo "$(echo 1 > /sys/class/gpio/gpio18/value)"
        echo "LED Acceso"
        sleep 0.5
    else
        echo "$(echo 0 > /sys/class/gpio/gpio18/value)"
        echo "LED Spento"
        sleep 0.5
    fi
done
```

Per prima cosa forniamo allo script i permessi:

```
# chmod +x PulsanteLed.sh
```

Poi lo eseguiamo:

```
# bash PulsanteLed.sh
```

E' possibile che se avete fatto qualche prova di accensione del LED oppure interrompete l'esecuzione dello script e lo riavviate, pur funzionando, il programma visualizzi un messaggio di errore con scritto "Dispositivo o risorsa occupata". E' perfettamente normale, significa che trova già impostato i GPIO 17 e 18 nella cartella "export". Per eliminare il messaggio basterà eseguire questo comando, ovviamente fornendo l'opportuno numero di GPIO:

```
# echo 18 > /sys/class/gpio/unexport
```

In questo modo la porta 18 verrà inserita nella cartella "unexport" e quindi non sarà più considerata come dispositivo funzionante. Se riuscite lo script i messaggi di errore dovrebbero sparire.

IMPORTANTE: questa possibilità offerta dai progettisti di RPi si adatta inoltre a qualsiasi linguaggio, se ci pensiamo bene basta infatti includere nel proprio codice il controllo dei comandi sopra descritti, usando le istruzioni per la riga di comando, per poter leggere e pilotare tutte le porte GPIO anche senza una libreria.

Facciamolo in PHP

Ho tenuto per ultimo il PHP sia perché è un linguaggio web sia perché è quello più lungo da descrivere. Chi conosce il PHP scoprirà che usarlo è comunque elementare. PHP non ha una libreria propria pertanto si dovrà utilizzare, come accennato sopra, programmi o utility esterne richiamandole con istruzioni del linguaggio che eseguono la riga di comando (es. `exec`, `system`, ecc.).

In questo caso io ho utilizzato l'utility GPIO che viene installata insieme a WiringPi (vedi paragrafo sul C). Questa utility ha pochissimi comandi e una semplicissima sintassi che vi sarà subito chiara quando analizzerete il codice.

Ma andiamo con ordine. Per prima cosa dobbiamo installare un web server e il PHP. Io preferisco Apache ma anche Nginx va benone.

```
# apt-get install apache2
```

Proviamo subito se funziona aprendo il browser e digitando l'IP del localhost:

```
http://127.0.0.1
```

Apparirà il testo "It's Work" riferito al fatto che Apache sta funzionando correttamente. Installiamo ora PHP e il modulo necessario per farlo girare su Apache.

```
# apt-get install php5 libapache2-mod-php5
```

Adesso, tutto il codice che scriviamo, dobbiamo metterlo in `/var/www/` cioè la cartella di default dove gli script PHP verranno eseguiti e che corrisponde al localhost. Volendo fare una prova se PHP gira correttamente portiamoci dentro la cartella `/var/www/` e creiamo con NANO un file

di nome `test.php`.

```
# nano test.php
```

Scriviamo la classica istruzione di informazioni su PHP e salviamo (CTRL-O poi Invio e infine CTRL-X)

```
<?php
phpinfo;
?>
```

Torniamo al browser e digitiamo `http://127.0.0.1/test.php` e comparirà la lunga tabella informativa su PHP a riprova che l'interprete funziona. Adesso creiamo una vera pagina HTML + PHP per accendere e spegnere il LED.

```

<?php
/* l'istruzione exec esegue un comando da shell */
exec("gpio read 0", $stato);
$led = '';

if (isset($_POST['accendi']))
{
    exec("gpio mode 1 out");
    exec("gpio write 1 1");
    $led = '';
}
if (isset($_POST['spegni']))
{
    exec("gpio mode 1 out");
    exec("gpio write 1 0");
    $led = '';
}
if (isset($_POST['lettura']))
{
    exec("gpio read 0", $stato);
}
?>

<html>
<head>
    <title> EMCelettronica - by Gabriele Guizzardi </title>
</head>

<body>
    <H2 align=center> Accensione LED tramite pulsante web </H2>

    <form method="post" action="">
        <p align=center>
            LED : <button name="accendi">Accendi</button>
            <button name="spegni">Spegni</button>
            <?php echo $led; ?>
        </p>
    </form>

    <form method="post" action="">
        <p align=center>
            Valore Pulsante : <button name="lettura"> Leggi stato </button>
            <input type="text" name="valor" value= "<?php echo $stato[0]; ?>" >
        </p>
        <p align=center>Tenete premuto il pulsante del circuito e cliccate Leggi Stato.</p>
    </form>

</body>
</html>

```

In pratica abbiamo una prima parte del codice che controllerà quale pulsante HTML verrà cliccato, in base a questa informazione eseguirà l'utility GPIO. Quindi con "gpio mode 1 out" indichiamo che vogliamo che il pin 12 (numerazione WiringPi, vedi tabella) sia settato in uscita (Led) per poi portarlo alto con "gpio write 1 1". In modo opposto "gpio write 1 0" spegnerà il led. Per leggere lo stato del pulsante useremo "gpio read 0" cioè leggi lo stato del pin 11 dove abbiamo collegato il pulsante. Se lo stato è 1 il pulsante è a riposo, se è 0 il pulsante è premuto. Per vedere la nostra pagina in azione basta tornare sul browser e digitare `http://127.0.0.1/PulsanteLed.php` ed avremo la seguente videata.

plice e pressoché immediato, Rpi vi da tutti gli strumenti che servono.

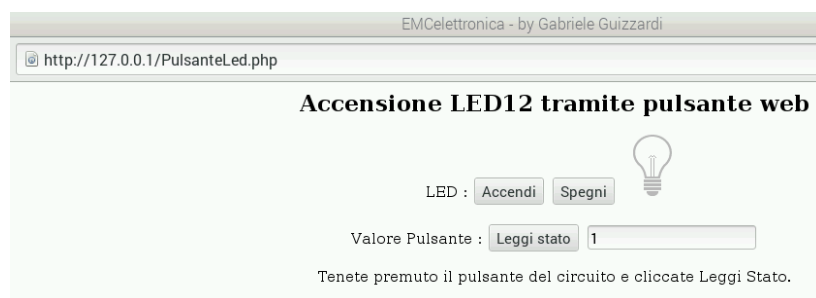
articolo1.tar.gz

Fonti e documentazione

Libreria wiringPi [<http://wiringpi.com/>]

Libreria pi4j [<http://pi4j.com/>]

PHP [<http://php.net/>]



Leggendo il codice noterete che ho utilizzato anche un paio di immagini, "acceso.png" e "spento.png", per rappresentare la lampadina.

Conclusioni

Programmare Raspberry Pi è un mondo che abbiamo solo scalfito, con le nuove versioni, poi, questo mondo è pure in espansione, entro fine anno infatti sulla RPi 2 sarà possibile far girare anche una versione di Windows 10. Da adesso quindi non avete più scuse, come avete potuto modo di vedere interagire con la GPIO è sem-

L'autore è a disposizione nei commenti per eventuali approfondimenti sul tema dell'Articolo.
Di seguito il link per accedere direttamente all'articolo sul Blog e partecipare alla discussione:
<http://it.emcelettronica.com/programmiamo-raspberry-pi-con-tutti-i-linguaggi>

Un altro portento della matematica: la trasformata di Laplace



di Marco Giancola
10 Settembre 2015

$$\int_0^{\infty} e^{-st} f(t) dt$$

*Dopo aver parlato, nei due precedenti articoli, della **trasformata di Fourier** e della **trasformata Z**, non potevamo non occuparci di un'altra importantissima trasformata ampiamente utilizzata in ambito scientifico e tecnologico: quella di Laplace. Benché venga chiamata così in onore di **Pierre Simon Laplace** (in quanto la utilizzò nei suoi studi riguardanti la teoria della probabilità), essa in realtà venne sviluppata da Eulero. Tra le sue numerose applicazioni, vi sono: la risoluzione di equazioni integro-differenziali, l'analisi dei circuiti e l'elaborazione dei segnali. In questo articolo, mostreremo in che modo con la trasformata di Laplace sia possibile, ad esempio, studiare gli effetti dei carichi sulle travi e determinare il comportamento di un circuito elettrico. Vedremo inoltre come essa, grazie al software MATLAB, possa essere calcolata facilmente anche da chi non possiede elevate competenze matematiche.*

1 DEFINIZIONE

Sia $f(t)$ una funzione a valori reali o complessi definita (almeno) sul semiasse reale non negativo ($t \geq 0$). Si dice che f è **trasforma-**

bile secondo Laplace (o **L-trasformabile**) se esiste un numero complesso s tale che la funzione (di t) $e^{-st}f(t)$ sia sommabile su $[0, +\infty)$, ossia tale che:

$$\int_0^{+\infty} |e^{-st} f(t)| dt < +\infty$$

ovvero, se ci riferiamo all'integrabilità secondo Lebesgue, tale che: $e^{-st}f(t) \in L^1([0, +\infty))$. In tal caso, l'integrale

$$\int_0^{+\infty} e^{-st} f(t) dt$$

è ben definito e prende il nome di **integrale di Laplace**. Se $e^{-st}f(t)$ è sommabile su $[0, +\infty)$ per un certo valore complesso s_0 , allora lo è per ogni $s \in \mathbb{C}$ tale che $\text{Re}(s) > \text{Re}(s_0)$; infatti:

e pertanto $e^{-st}f(t)$, essendo maggiorata in modulo da una funzione sommabile su $[0, +\infty)$, è anch'essa sommabile su $[0, +\infty)$.

Da ciò si evince che l'insieme dei valori di s che rendono sommabile $e^{-st}f(t)$ su $[0, +\infty)$, se non è vuoto, è costituito da tutti i numeri complessi la cui parte reale è maggiore di un valore che viene indicato con $\sigma[f]$ e prende il nome di **ascissa di convergenza** di f . I punti corrispondenti a tali numeri complessi sono ovviamente quelli del semipiano destro individuato dalla retta $x = \sigma[f]$, detti rispettivamente **semipiano di convergenza** e **retta di convergenza**.

L'integrale di Laplace di una funzione L-trasformabile $f(t)$ è dunque una funzione definita in $\{s \in \mathbb{C} \mid \text{Re}(s) > \sigma[f]\}$; essa viene chiamata trasformata di Laplace e si indica con

$$\mathcal{L}[f](s) = F(s) := \int_0^{+\infty} e^{-st} f(t) dt.$$

È evidente che la trasformata di Laplace rientra nella categoria delle **trasformate integrali**. Il suo

nucleo integrale è ovviamente e^{-st} .

1.1 LA TRASFORMATATA BILATERA DI LAPLACE

Se $f(t)$ è definita su tutto $\mathbb{R}$ ed è tale che $e^{-st}f(t) \in L^1(\mathbb{R})$ per certi valori di s , è possibile definire la **trasformata bilatera di Laplace** di $f(t)$:

$$F(s) = \int_{-\infty}^{+\infty} e^{-st} f(t) dt$$

Per determinare il dominio di $F(s)$ (ovvero il sottoinsieme di $\mathbb{C}$ che rende sommabile $e^{-st}f(t)$ su $\mathbb{R}$), consideriamo le seguenti funzioni definite solo per $t > 0$:

$$\begin{aligned} f_+(t) &= \text{restrizione di } f(t) \text{ a } t > 0, \\ f_-(t) &= \text{restrizione di } f(-t) \text{ a } t > 0 \end{aligned}$$

e indichiamo con $F_+(s)$ e $F_-(s)$ le loro rispettive trasformate di Laplace (unilatera). Si ha:

$$\begin{aligned} F(s) &= \int_{-\infty}^{+\infty} e^{-st} f(t) dt = \int_{-\infty}^0 e^{-st} f(t) dt + \int_0^{+\infty} e^{-st} f(t) dt \\ &= \int_0^{+\infty} e^{st} f(-t) dt + \int_0^{+\infty} e^{-st} f(t) dt \\ &= \int_0^{+\infty} e^{st} f_-(t) dt + \int_0^{+\infty} e^{-st} f_+(t) dt = F_-(-s) + F_+(s) \end{aligned}$$

Pertanto, il dominio di $F(s)$ è l'intersezione dei domini di $F_+(s)$ e $F_-(-s)$:

$$\begin{aligned} &\{s \in \mathbb{C} \mid \text{Re}(s) > \sigma[f_+]\} \cap \{s \in \mathbb{C} \mid \text{Re}(-s) > \sigma[f_-]\} \\ &= \{s \in \mathbb{C} \mid \text{Re}(s) > \sigma[f_+]\} \cap \{s \in \mathbb{C} \mid \text{Re}(s) < -\sigma[f_-]\} \\ &= \{s \in \mathbb{C} \mid \sigma[f_+] < \text{Re}(s) < -\sigma[f_-]\} \end{aligned}$$

e quindi la trasformata bilatera di f esiste soltanto se $\sigma[f_+] < -\sigma[f_-]$.

Se il dominio della trasformata bilatera di Laplace di una funzione f include l'asse imma-

ginario (l'asse delle ordinate), **f è dotata anche della trasformata di Fourier, e su tale asse le due trasformate coincidono; infatti, per $s = i\lambda$ la trasformata bilatera di Laplace diventa quella di Fourier.** In altre parole, la prima generalizza la seconda.

2 ESEMPI

2.1 1° ESEMPIO

Consideriamo la funzione di **Heaviside**:

$$H(t) = \begin{cases} 1 & \text{se } t \geq 0 \\ 0 & \text{se } t < 0 \end{cases}$$

Poiché

$$|e^{-st}| = e^{-\operatorname{Re}(s)t}$$

e^{-st} risulta sommabile su $[0, +\infty)$ se e solo se $\operatorname{Re}(s) > 0$, ovvero, per s tale che $\operatorname{Re}(s) > 0$, $e^{-st}H(t)$ è sommabile su tutto l'asse reale. Dunque $\sigma[f] = 0$. Le trasformate di Laplace, unilatera e bilatera, di $H(t)$ coincidono e sono uguali a:

$$\begin{aligned} F(s) &= \int_0^{+\infty} e^{-st} dt = \lim_{x \rightarrow +\infty} \frac{1}{-s} \int_0^x e^{-st} d(-st) = \lim_{x \rightarrow +\infty} \frac{1}{-s} [e^{-st}]_0^x \\ &= \lim_{x \rightarrow +\infty} \frac{1}{-s} (e^{-sx} - 1) = \frac{1}{s} \left(1 - \lim_{x \rightarrow +\infty} e^{-\operatorname{Re}(s)x - i\operatorname{Im}(s)x} \right) \\ &= \frac{1}{s} \left(1 - \lim_{x \rightarrow +\infty} e^{-\operatorname{Re}(s)x} \cos(\operatorname{Im}(s)x) \right. \\ &\quad \left. + i \lim_{x \rightarrow +\infty} e^{-\operatorname{Re}(s)x} \sin(\operatorname{Im}(s)x) \right) = \frac{1}{s} (1 - 0 + 0i) = \frac{1}{s} \end{aligned}$$

$H(t)$ è un esempio di funzione che ammette la trasformata di Laplace ma non quella di Fourier, non essendo di classe $L^1(\mathbb{R})$. Difatti, il dominio di $F(s)$ non include l'asse immaginario.

2.2 2° ESEMPIO

Consideriamo la funzione $f(t) = e^{at}$ con $a \in \mathbb{C}$. La funzione $e^{-st}e^{at} = e^{-(s-a)t}$ è sommabile per $t \in [0, +\infty)$ se e solo se $\operatorname{Re}(s-a) = \operatorname{Re}(s) - \operatorname{Re}(a) > 0$, ossia $\operatorname{Re}(s) > \operatorname{Re}(a)$ e per-

tanto $\sigma[f] = \operatorname{Re}(a)$. Il calcolo della trasformata di Laplace di $f(t)$ è simile a quello dell'esempio precedente:

$$\mathcal{L}[f](s) = \int_0^{+\infty} e^{-(s-a)t} dt = \frac{1}{s-a}.$$

Per $a = 0$, si ritrova il risultato relativo alla funzione di Heaviside.

2.3 3° ESEMPIO

Calcoliamo la trasformata di Laplace della funzione caratteristica di un intervallo $[0, h)$:

$$\mathcal{X}_{[0,h)}(t) = H(t) - H(t-h) = \begin{cases} 1 & 0 \leq t < h \\ 0 & \text{altrove.} \end{cases}$$

Il prodotto di tale funzione per e^{-st} è chiaramente sommabile su tutto $\mathbb{R}$ per qualunque valore di s , e quindi $\sigma[f] = -\infty$. Per $s = 0$, l'integrale di Laplace si riduce a $\int dt$ calcolato tra 0 e h , e pertanto $F(0) = h$. Per $s \neq 0$, invece, si ha:

$$\mathcal{L}[\mathcal{X}_{[0,h)}](s) = \int_0^{+\infty} e^{-st} \mathcal{X}_{[0,h)}(t) dt = \int_0^h e^{-st} dt = \left[\frac{e^{-st}}{s} \right]_{t=0}^{t=h} = \frac{1 - e^{-sh}}{s}.$$

C'è dunque una **singolarità** in $s = 0$, che però risulta eliminabile:

$$\lim_{s \rightarrow 0} \mathcal{L}[\mathcal{X}_{[0,h)}](s) = \lim_{s \rightarrow 0} \frac{1 - e^{-sh}}{s} = \lim_{s \rightarrow 0} \frac{1 - e^{-sh}}{sh} h = h = \mathcal{L}[\mathcal{X}_{[0,h)}](0).$$

2.4 LE TRASFORMATE DI ALCUNE FUNZIONI NOTEVOLI

- **Delta di Dirac:**

$$\mathcal{L}\{\delta(t)\} = 1$$

- **Funzione gradino di Heaviside:**

$$\mathcal{L}\{\Theta(t)\} = \frac{1}{s}$$

- **Funzione esponenziale:**

$$\mathcal{L}\{e^{\alpha t}\} = \frac{1}{s - \alpha}$$

- **Seno:**

$$\mathcal{L}\{\sin(\alpha t)\} = \frac{\alpha}{s^2 + \alpha^2}$$

$$\mathcal{L}\{e^{\alpha t} \sin(\beta t)\} = \frac{\beta}{(s - \alpha)^2 + \beta^2}$$

- **Coseno:**

$$\mathcal{L}\{\cos(\alpha t)\} = \frac{s}{s^2 + \alpha^2}$$

$$\mathcal{L}\{e^{\alpha t} \cos(\beta t)\} = \frac{s - \alpha}{(s - \alpha)^2 + \beta^2}$$

- **Seno iperbolico:**

$$\mathcal{L}\{\sinh(\alpha t)\} = \frac{\alpha}{s^2 - \alpha^2}$$

- **Coseno iperbolico:**

$$\mathcal{L}\{\cosh(\alpha t)\} = \frac{s}{s^2 - \alpha^2}$$

3 ALCUNE PROPRIETÀ

3.1 UNICITÀ

Se due funzioni L-trasformabili $f(t)$ e $g(t)$ hanno la medesima trasformata di Laplace, esse coin-

3.2 LINEARITÀ

La trasformata di Laplace è lineare, ossia, per ogni coppia di funzioni L-trasformabili f_1 e f_2 con ascisse di convergenza $\sigma[f_1]$ e $\sigma[f_2]$ e per ogni coppia di numeri reali o complessi c_1 e c_2 , si ha:

$$\mathcal{L}[c_1 f_1 + c_2 f_2](s) = c_1 \mathcal{L}[f_1](s) + c_2 \mathcal{L}[f_2](s),$$

$$\forall s: \operatorname{Re}(s) > \max\{\sigma[f_1], \sigma[f_2]\}.$$

3.3 LIMITATEZZA

Se f è una funzione L-trasformabile con ascissa di convergenza $\sigma[f]$, per ogni $\sigma_0 > \sigma[f]$, la sua trasformata di Laplace $F(s)$ è limitata nel semipiano chiuso $\operatorname{Re}(s) \geq \sigma_0$. Inoltre, per ogni successione $\{s_n\}$ tale che $\sigma_0 \leq \operatorname{Re}(s_n) \rightarrow +\infty$ per $n \rightarrow \infty$, $F(s_n) \rightarrow 0$ per $n \rightarrow \infty$.

3.4 DERIVABILITÀ

Se f è una funzione L-trasformabile con ascissa di convergenza $\sigma[f]$, la sua trasformata di Laplace $F(s)$ è **olomorfa** nel semipiano $\operatorname{Re}(s) > \sigma[f]$. Inoltre, la funzione $-tf(t)$ è anch'essa L-trasformabile con ascissa di convergenza $\sigma[f]$ e si ha:

$$F'(s) = \mathcal{L}[-tf(t)](s),$$

e più in generale:

$$F^{(n)}(s) = \mathcal{L}[(-t)^n f(t)](s) = (-1)^n \mathcal{L}[t^n f(t)] \quad \text{di Laplace è:} \quad \text{Re}(s) > \sigma[f].$$

Da ciò si può ricavare che, se anche $f(t)/t$ è L-trasformabile, si ha:

$$\mathcal{L}\left[\frac{f(t)}{t}\right](s) = \int_s^\infty \mathcal{L}[f](\tau) d\tau \quad \text{Re}(s) > \sigma[f]$$

3.5 TRASFORMATA DI LAPLACE DELLE DERIVATE

Se f è una funzione L-trasformabile, continua e derivabile in $[0, +\infty)$ con derivata prima continua a tratti ed anch'essa L-trasformabile, per ogni s tale che $\text{Re}(s) > \max\{\sigma[f], \sigma[f']\}$, si ha:

$$\mathcal{L}[f'](s) = s \mathcal{L}[f](s) - f(0).$$

In generale, se la suddetta funzione f è di classe $C^{n-1}([0, +\infty))$ e $f^{(n)}$ verifica le ipotesi di cui sopra, vale la seguente formula:

$$\mathcal{L}[f^{(n)}](s) = s^n \mathcal{L}[f](s) - \left(s^{n-1} f(0) + s^{n-2} f'(0) + \dots + f^{(n-1)}(0) \right)$$

3.6 CONVOLUZIONE

Se f e g sono due funzioni L-trasformabili nulle per $t < 0$, il loro prodotto di convoluzione $f * g$ è L-trasformabile nel semipiano $\text{Re}(s) > \max\{\sigma[f], \sigma[g]\}$ e si ha:

$$\mathcal{L}[f * g](s) = \mathcal{L}[f](s) \cdot \mathcal{L}[g](s).$$

Se g è la funzione di Heaviside, il prodotto di convoluzione diventa:

$$(f * H)(t) = (H * f)(t) = \int_0^t f(\tau) d\tau$$

che è la primitiva di f nulla nell'origine. In virtù della precedente relazione, la sua trasformata

$$\mathcal{L}[H * f](s) = \frac{F(s)}{s}, \quad \text{Re}(s) > \max\{0, \sigma[f]\}$$

Si noti che tale risultato concorda con la formula del sottoparagrafo 3.5 e fornisce un'informazione supplementare: è sufficiente richiedere la L-trasformabilità della derivata di una funzione per ottenere automaticamente la L-trasformabilità della funzione.

3.7 ALTRE PROPRIETÀ

Se f è una funzione L-trasformabile nulla per $t < 0$, valgono le seguenti relazioni:

- $\mathcal{L}[f(ct)](s) = \frac{1}{c} \mathcal{L}[f(t)]\left(\frac{s}{c}\right) \quad \forall c > 0 : \operatorname{Re}(s) > c\sigma[f];$
- $\mathcal{L}[f(t - t_0)](s) = e^{-t_0 s} \mathcal{L}[f(t)](s) \quad \forall t_0 > 0 : \operatorname{Re}(s) > \sigma[f];$
- $\mathcal{L}[e^{at} f(t)](s) = \mathcal{L}[f(t)](s - a) \quad \forall a \in \mathbb{C} : \operatorname{Re}(s) > \sigma[f] + \operatorname{Re}(a).$

4 INVERSIONE DELLA TRASFORMATA DI LAPLACE

Occupiamoci ora del problema dell'inversione della trasformata di Laplace. Enunciamo un teorema di cui omettiamo la dimostrazione.

Sia $F(s)$ una **funzione analitica** definita nel semipiano $\{s \in \mathbb{C} \mid \operatorname{Re}(s) > \sigma_0\}$ e tale che $|F(s)| = O(1/|s|^k)$ per $|s| \rightarrow +\infty$, con $k > 1$. Allora, per ogni $\alpha > \sigma_0$, la funzione

$$f(t) = \frac{1}{2\pi i} \int_{\alpha - i\infty}^{\alpha + i\infty} e^{st} F(s) ds$$

è L-trasformabile, nulla per $t < 0$, continua su $\mathbb{R}$, indipendente da α ed avente la F come trasformata di Laplace. Essa prende il nome di **trasformata inversa o antitrasformata di Laplace**, mentre la relazione che la definisce si chiama **formula di Riemann-Fourier o di Bromwich-Mellin**.

Per il **teorema dei residui**, si ha:

$$f(t) = \frac{1}{2\pi i} \int_{\alpha - i\infty}^{\alpha + i\infty} e^{st} F(s) ds = \sum_{j=1}^n \operatorname{res}(e^{st} F(s), s_j),$$

dove gli s_j sono i **punti singolari** della funzione $e^{st} F(s)$.

5 CALCOLARE TRASFORMATA E ANTITRASFORMATA DI LAPLACE CON MATLAB

Il software MATLAB rende alquanto agevole il calcolo della trasformata e dell'antitrasformata di Laplace. Vediamo un

paio di esempi.

Supponiamo di voler calcolare la trasformata di Laplace della funzione

$$f(t) = -1.25 + 3.5te^{-2t} + 1.25e^{-2t}$$

Ecco cosa va digitato nella finestra dei comandi e l'output che appare:

```
>> syms t s
>> f=-1.25+3.5*t*exp(-2*t)+1.25*exp(-2*t);
>> F=laplace(f,t,s)

F =

-5/4/s+7/2/(s+2)^2+5/4/(s+2)

>> simplify(F)

ans =

(s-5)/s/(s+2)^2

>> pretty(ans)
```

$$\frac{s - 5}{s (s + 2)^2}$$

Come si può vedere, occorre:

1. dichiarare le variabili simboliche t e s con il comando `syms`;
2. definire la funzione $f(t)$;
3. **calcolare la trasformata di Laplace tramite la funzione `laplace`.**

È possibile rendere l'output più comprensibile mediante le funzioni `simplify` e `pretty`. Come si

legge chiaramente, il risultato è:

$$F(s) = \frac{(s-5)}{s(s+2)^2}$$

In alternativa, si può scrivere la funzione $f(t)$ direttamente all'interno delle parentesi della funzione *laplace*:

```
>> F2=laplace(-1.25+3.5*t*exp(-2*t)+1.25*exp(-2*t))
```

In maniera perfettamente analoga, **utilizzando la funzione *ilaplace*, è possibile calcolare la trasformata inversa**. Vediamo come riottenere la funzione $f(t)$ dell'esempio precedente a partire dalla sua trasformata $F(s)$:

```
>> syms t s
>> F=(s-5)/(s*(s+2)^2);
>> ilaplace(F)
ans =
-5/4+(7/2*t+5/4)*exp(-2*t)
>> simplify(ans)
ans =
-5/4+7/2*t*exp(-2*t)+5/4*exp(-2*t)
>> pretty(ans)
- 5/4 + 7/2 t exp(-2 t) + 5/4 exp(-2 t)
```

che è proprio $f(t)$. In alternativa, si può scrivere:

```
>> ilaplace((s-5)/(s*(s+2)^2))
```

6 ALCUNE APPLICAZIONI

6.1 RISOLUZIONE DI EQUAZIONI DIFFERENZIALI LINEARI A COEFFICIENTI COSTANTI

Consideriamo il **problema di Cauchy** relativo ad un'equazione differenziale lineare del secondo ordine, a coefficienti costanti, non omogenea:

$$\begin{cases} y''(t) + ay'(t) + by(t) = g(t) & t \geq 0 \\ y(0) = y_0 \quad y'(0) = y_1, \end{cases}$$

con $g(t)$ funzione L-trasformabile. Applicando la trasformata di Laplace all'equazione e sfruttan-

do la formula del sottoparagrafo 3.5, si ha:

$$s^2Y(s) - sy_0 - y_1 + a(sY(s) - y_0) + bY(s) = \mathcal{L}[g](s),$$

ossia:

$$(s^2 + as + b)Y(s) = y_0(s + a) + y_1 + \mathcal{L}[g](s),$$

da cui si ricava:

$$Y(s) = \frac{y_0(s + a) + y_1}{s^2 + as + b} + \frac{\mathcal{L}[g](s)}{s^2 + as + b} = Y_1(s) + Y_2(s).$$

A questo punto, rimane solo da antitrasformare $Y(s)$ per ottenere la soluzione $y(t)$ del problema di Cauchy.

Se l'equazione differenziale è omogenea ($g(t) = 0$), la trasformata di Laplace della soluzione è $Y(s) = Y_1(s)$, dunque una funzione razionale; questa circostanza si verifica per ogni equazione differenziale lineare a coefficienti costanti omogenea, indipendentemente dall'ordine. In questo caso, $Y(s)$ è spesso facilmente antitrasformabile mediante la **scomposizione in fratti semplici** e la conoscenza delle trasformate delle più comuni funzioni notevoli (vedi sottoparagrafo 2.4).

6.1.1 1° ESEMPIO

Consideriamo il seguente problema di Cauchy:

$$\begin{cases} y''(t) + 4y'(t) + 3y(t) = 0 & t \geq 0 \\ y(0) = 0 \quad y'(0) = 1. \end{cases}$$

Trasformando ambo i membri si trova

$$(s^2 + 4s + 3)Y(s) = 1 \iff Y(s) = \frac{1}{s^2 + 4s + 3} = \frac{1}{(s+1)(s+3)}.$$

Decomponendo $Y(s)$ in fratti semplici, si ha:

$$\frac{A}{s+1} + \frac{B}{s+3} = \frac{As + 3A + Bs + B}{(s+1)(s+3)} = \frac{1}{(s+1)(s+3)}$$

$$\Rightarrow \begin{cases} A + B = 0 \\ 3A + B = 1 \end{cases} \Rightarrow \begin{cases} A = 1/2 \\ B = -1/2 \end{cases} \Rightarrow Y(s) = \frac{1/2}{s+1} - \frac{1/2}{s+3}$$

Poiché la trasformata di e^{at} è $1/(s-a)$ (vedi sottoparagrafo 2.2), ovvero l'antitrasformata di $1/(s-a)$ è e^{at} , possiamo concludere che la soluzione del problema è:

$$y(t) = \frac{e^{-t}}{2} - \frac{e^{-3t}}{2}.$$

6.1.2 2° ESEMPIO

Vogliamo risolvere, tramite la trasformata di Laplace, il seguente problema di Cauchy:

$$\begin{cases} y''(t) + y(t) = te^t \\ y(0) = 1, y'(0) = 1. \end{cases}$$

Trasformiamo entrambi i membri dell'equazione:

$$Y(s) = \frac{s+1}{s^2+1} + \frac{1}{(s^2+1)(s-1)^2} = Y_1(s) + Y_2(s).$$

Per $Y_1(s)$ si ha

$$Y_1(s) = \frac{s}{s^2+1} + \frac{1}{s^2+1} \implies y_1(t) = \cos(t) + \sin(t).$$

Mentre per $Y_2(s)$:

$$Y_2(s) = \mathcal{L}[te^t] \mathcal{L}[\sin(t)] \implies y_2(t) = te^t * \sin(t).$$

6.1.3 3° ESEMPIO

Consideriamo il seguente problema di Cauchy:

$$\begin{cases} y''(t) + y(t) = 3 \\ y(5) = 0 \quad y'(5) = 1. \end{cases}$$

Stavolta, le condizioni iniziali non sono assegnate in $t = 0$. Tuttavia, ponendo:

$$\bar{y}(t) = y(t+5)$$

il problema di Cauchy diventa:

$$\begin{cases} \bar{y}''(t) + \bar{y}(t) = 3 \\ \bar{y}(0) = 0 \quad \bar{y}'(0) = 1. \end{cases}$$

A questo punto, possiamo procedere come negli esempi precedenti: trasformiamo entrambi i membri dell'equazione, ricaviamo $Y(s)$ e la scomponiamo in fratti semplici:

$$Y(s) = \frac{3}{s(s^2 + 1)} + \frac{1}{s^2 + 1} = \frac{a}{s} + \frac{b}{s - i} + \frac{c}{s + i} + \frac{1}{s^2 + 1}, \quad a, b, c \in \mathbb{R}.$$

Risolvendo si trova $a = 3$, $b = -\frac{3}{2}$, $c = -\frac{3}{2}$. Si riesce dunque facilmente ad antitrasformare:

$$\bar{y}(t) = \sin(t) + 3 - \frac{3}{2}(e^{-it} + e^{it}) = \sin(t) + 3(1 - \cos(t)).$$

Pertanto, la soluzione del problema originario è:

ESEMPIO Risolvere l'equazione integrale:

$$y(t) = \sin(t - 5) + 3(1 - \cos(t - 5)). \quad \varphi(t) = t + \int_0^t \sin(t - s) \varphi(s) ds.$$

6.2 RISOLUZIONE DI EQUAZIONI INTEGRALI DI TIPO CONVOLUTIVO

La trasformata di Laplace trova applicazione anche nella risoluzione di equazioni integrali di tipo convolutivo, ossia relazioni del tipo:

$$f(t) = g(t) * f(t) + u(t)$$

chiamate così per via della presenza del prodotto di convoluzione $g(t) * f(t)$. Infatti, se f , g e u sono L-trasformabili, applicando la trasformata di Laplace ad ambo i membri dell'equazione, ottengo:

$$F(s) = G(s)F(s) + U(s)$$

da cui:

$$F(s) = U(s)/(1 - G(s))$$

Quindi, se $1/(1 - G(s))$ risulta essere antitrasformabile, indicando con $h(t)$ la sua antitrasformata, otteniamo la soluzione

$$f(t) = h(t) * u(t)$$

Soluzione Essendo $\mathcal{L}\{t\}(z) = 1/z^2$, $\mathcal{L}\{\sin t\}(z) = 1/(1 + z^2)$, l'equazione si trasforma nella:

$$\Phi(z) = z^{-2} + (1 + z^2)^{-1} \Phi(z),$$

ove si è posto $\Phi = \mathcal{L}\{\varphi\}$. Dalla precedente segue:

$$\Phi(z) = \frac{1}{z^2} + \frac{1}{z^4},$$

e quindi la soluzione è $\varphi(t) = t + t^3/3!$.

6.3 STUDIO DEGLI EFFETTI DEI CARICHI SULLE TRAVI

Consideriamo una trave orizzontale di lunghezza l e modellizziamola con il segmento $[0, l]$ dell'asse x . Supponiamo che su di essa agisca un carico verticale $F(x)$, con $x \in [0, l]$. In tali condizioni, la trave subisce uno spostamento trasversale $y(x)$ che verifica la seguente equazione differenziale del 4° ordine:

$$y^{IV}(x) = kF(x), \quad 0 < x < l,$$

dove k è una costante che dipende dalle caratteristiche della trave. Per calcolare $y(x)$, occorre risolvere il [problema ai limiti](#) ottenuto associando all'equazione precedente delle condizioni agli estremi 0 e l che dipendono da come la trave è vincolata a tali estremi. Si distinguono 3 casi:

- a) Estremità incastrata : $y = y' = 0$.
- b) Estremità appoggiata : $y = y'' = 0$.
- c) Estremità libera : $y'' = y''' = 0$.

Proviamo a risolvere il problema nel caso in cui la trave sia appoggiata alle due estremità, e supponendo F costante (distribuzione uniforme del carico), ossia:

$$yIV = kF, 0 < x < l$$

$$y(0) = y''(0) = 0$$

$$y(l) = y''(l) = 0$$

Applicando la trasformata di Laplace all'equazione differenziale e ponendo $c_1 = y'(0)$ e $c_2 = y'''(0)$, otteniamo:

$$z^4 Y(z) - c_1 z^3 - c_2 = \frac{kF}{z^4},$$

da cui segue:

$$Y(z) = \frac{c_1}{z^2} + \frac{c_2}{z^4} + \frac{kF}{z^8},$$

e quindi antitrasformando:

$$y(x) = c_1 x + c_2 \frac{x^3}{3!} + \frac{kF}{4!} x^4.$$

Le condizioni $y(l) = y''(l) = 0$ ci consentono di determinare le due costanti:

$$c_1 = \frac{kF}{24} \ell^3, \quad c_2 = -\frac{kF}{2} \ell,$$

e pertanto lo spostamento trasversale della trave è:

$$y(x) = \frac{kF}{24} x(\ell - x)(\ell^2 + \ell x - x^2).$$

6.4 ANALISI DEI CIRCUITI ELETTRICI

Il funzionamento di un circuito elettrico è descritto da un'equazione differenziale in cui la funzione incognita è la corrente $i(t)$ o la tensione $v(t)$. Per tale motivo, l'analisi di un circuito elettrico include la risoluzione dell'equazione differenziale ad esso associata.

Ad esempio, un circuito RCL è modellizzato matematicamente da un'equazione differenziale del secondo ordine, a coefficienti costanti, di questo tipo:

$$Li''(t) + Ri'(t) + i(t)/C = v(t)$$

dove R , C e L sono rispettivamente i valori di resistenza, capacità e induttanza.

Nello studio dei circuiti RCL non ha grande interesse il comportamento per tempi piccoli, bensì quello a regime, per tempi grandi. Da un punto di vista matematico, ciò significa che non hanno particolare rilevanza le condizioni iniziali del problema di Cauchy per l'equazione precedente. Possiamo quindi supporre, per semplicità, $i(0) = i'(0) = 0$. Procedendo allora come negli esempi del sottoparagrafo 6.1, si ha:

$$I(s) = V(s)/(Ls^2 + Rs + 1/C)$$

Ponendo:

$$H(s) = 1/(Ls^2 + Rs + 1/C)$$

si ottiene:

$$I(s) = H(s)V(s)$$

e quindi:

$$i(t) = (\mathcal{L}^{-1}[H] * v)(t)$$

Pertanto, una volta nota $H(s)$, siamo in grado di determinare il comportamento del circuito per ogni $v(t)$ L-trasformabile. $H(s)$ è quindi una caratteristica intrinseca del circuito e prende il nome di **funzione di trasferimento**.

6.5 I SISTEMI LTI TEMPO-CONTINUI E LA FUNZIONE DI TRASFERIMENTO

Nel [precedente articolo](#), avevamo definito, mediante la trasformata Z, la funzione di trasferimento $H(z)$ di un [sistema LTI](#) tempo-discreto. In maniera perfettamente analoga, **tramite la trasformata di Laplace, si definisce la funzione di trasferimento $H(s)$ di un sistema LTI tempo-continuo** (di cui abbiamo già visto un esempio nel sottoparagrafo precedente): **è la trasformata di Laplace della risposta impulsiva $h(t)$ del suddetto sistema, ed è pertanto la funzione di rete che esprime la relazione algebrica tra il segnale in ingresso $x(t)$ e quello in uscita $y(t) = h(t)*x(t)$ nel dominio di s (detta anche frequenza complessa), descrivendo in tal modo il comportamento del sistema.**

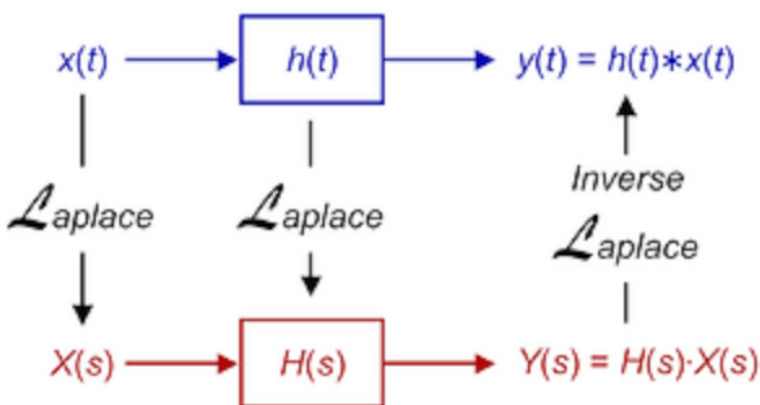
Dunque, **per calcolare il segnale di uscita $y(t)$, è sufficiente antitrasformare il prodotto $H(s)X(s)$ delle trasformate di Laplace di $h(t)$ e $x(t)$. Altrimenti, se non conosco $h(t)$, per determinare $y(t)$, posso sempre risolvere l'equazione differenziale lineare a coefficienti costanti che descrive il sistema LTI. In entrambi i casi, utilizzo la trasformata di Laplace.**

Si può dimostrare che **un sistema LTI è stabile** (ossia, ad ogni ingresso $x(t)$ limitato corrisponde un'uscita $y(t)$ limitata) **se tutti i poli della sua funzione di trasferimento $H(s)$ hanno parte reale negativa**, ovvero si trovano a sinistra dell'asse immaginario. Ai fini della stabilità del sistema, nessun particolare vincolo si pone, invece, sugli zeri della funzione di trasferimento. Nel caso di funzioni $H(s)$ rappresentative di sistemi descritti da equazioni differenziali con coefficienti reali (esprese quindi da rapporti di polinomi in s con coefficienti reali), i poli e gli zeri sono reali o, a coppie, complessi coniugati. Si consideri, ad esempio, la funzione di trasferimento:

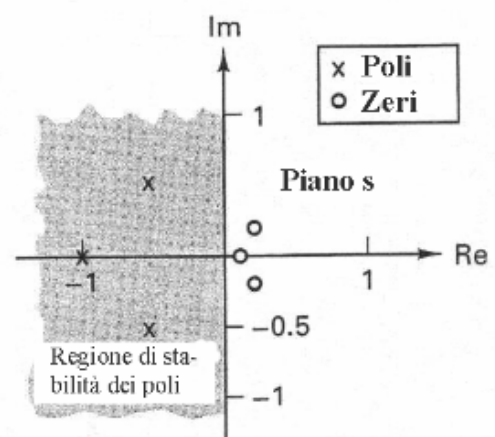
$$H(s) = \frac{s^3 - 0.5s^2 + 0.12s - 0.008}{s^3 + 2s^2 + 1.5s + 0.5}.$$

La distribuzione dei suoi poli e dei suoi zeri è riportata nella seguente figura:

Time domain



Frequency domain



Poli: $s = -1, -0,5 \pm 0,5i$.

Zeri: $s = 0, 1, 0,2 \pm 0,2i$.

Il sistema è certamente stabile.

Dopo quanto si è detto in questo sottoparagrafo, dovrebbe risultare chiara **la corrispondenza tra la trasformata di Laplace e la trasformata Z: la seconda può essere vista come l'equivalente discreto della prima.**

Come avevamo detto nel precedente articolo, un esempio di sistemi LTI tempo-discreti è rappresentato dai **filtri numerici**. Si intuisce facilmente che **i filtri analogici costituiscono invece un esempio di sistemi LTI tempo-continui.**

La trasformata di Laplace

Laplace Transforms with MATLAB

Metodi Matematici per l'Ingegneria

Pierre Simon Laplace

Trasformata di Laplace

TRASFORMATA DI LAPLACE

Trasformata inversa di Laplace

Trasformate Integrali di Fourier e di Laplace

BIBLIOGRAFIA

Gilardi G.: *analisi tre*; Milano, McGraw-Hill Libri Italia srl, 1994.

Stallo C.: *MODULO DI ELABORAZIONE DEI SEGNALI*, Master universitario in “Lo spazio: Galileo, telecomunicazioni e formazione – SPA”; Roma, Scuola IaD – Università Tor Vergata, 2007.

BIBLIOGRAFIA ON-LINE

APPLICAZIONI DELLA TRASFORMATA DI LAPLACE

Funzione di trasferimento

Inversione della trasformata di Laplace

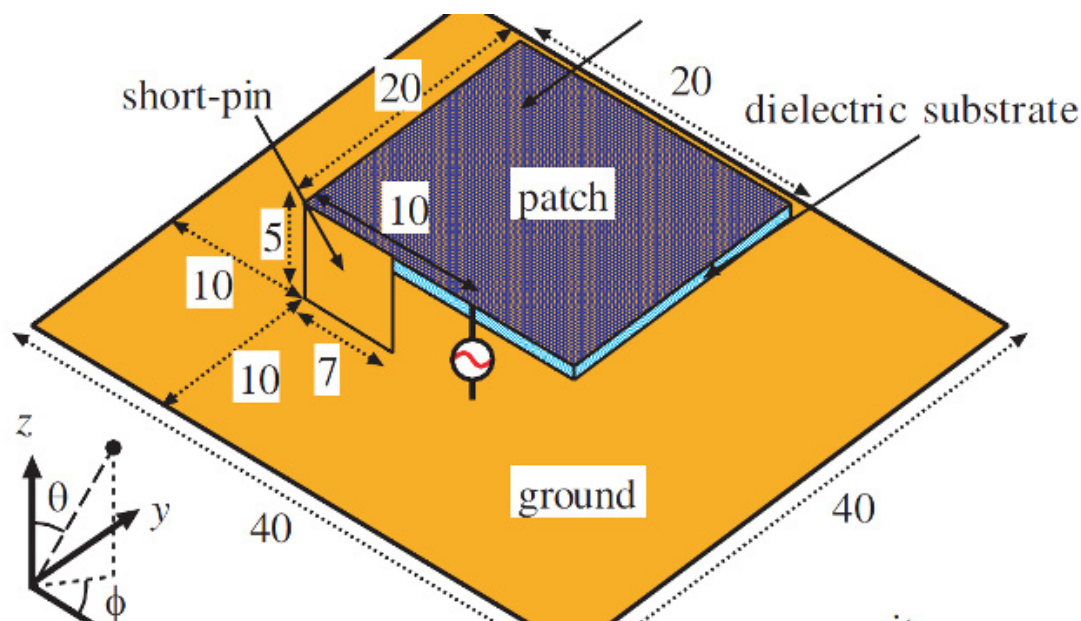
La trasformata di Laplace

L'autore è a disposizione nei commenti per eventuali approfondimenti sul tema dell'Articolo.
Di seguito il link per accedere direttamente all'articolo sul Blog e partecipare alla discussione:
<http://it.emcelettronica.com/un-altro-portento-della-matematica-la-trasformata-di-laplace>

Antenne trasparenti: la tecnologia del mimetismo



di XectX
15 Settembre 2015



Nel corso degli ultimi decenni la corsa all'integrazione si è manifestata nella ricerca di soluzioni di progetto volte alla continua riduzione delle dimensioni dei dispositivi elettronici, parallelamente alla necessità di prestazioni sempre maggiori. Nelle nostre mani sono arrivati dispositivi sempre più piccoli e portatili, frutto di rivoluzioni nell'ambito della progettazione, dei processi e dei materiali utilizzati. In questo articolo esploriamo le proprietà di un materiale molto particolare e vediamo la sua innovativa applicazione nell'ambito delle antenne trasparenti.

Come ogni progetto, quello di un antenna è vincolato a specifiche tipiche dell'ambito in cui si opera e caratteristiche della particolare applicazione. L'antenna deve avere le dimensioni giuste per risuonare alle frequenze di interesse, pertanto un modello di antenna può rivelarsi una scelta migliore di un'altra a seconda del contesto e dello spazio che si può occupare (con riferimento, nella fattispecie, alle antenne

stampate su una PCB). Tra le varie alternative, già a partire dalla fine del secolo scorso, si è rivelata attuabile l'idea di riutilizzare lo spazio a disposizione. Come? Un esempio è stampare antenne trasparenti sui display dei dispositivi o delle strumentazioni (si pensi ad esempio ad apparecchiature biomedicali in grado di comunicare tra loro), oppure applicando le medesime antenne sulle celle solari di piccoli satelliti,

dei quali si semplificherebbe in parte il progetto, con le relative conseguenze. La possibilità di un'antenna trasparente deriva direttamente dall'utilizzo di un materiale che conservi buone proprietà elettriche e, allo stesso tempo, buone proprietà di trasparenza ottica.

Passiamo dunque ad indagare le caratteristiche dei materiali conduttivi, per capire se e come sia possibile ottenere tale risultato.

IL LEGAME TRA LA CONDUCIBILITÀ E LA TRASPARENZA

Per la descrizione dei fenomeni elettrici all'interno di un materiale conduttivo, risulta conveniente l'utilizzo del modello di Drude. Consideriamo un materiale omogeneo ed isotropo (in cui i fenomeni elettrici non hanno direzioni preferenziali), sotto l'azione di un campo elettrico $\mathbf{E}$, possiamo assumere che sugli elettroni di conduzione, non vincolati ad alcun nucleo atomico, agiscano una forza elettrica ed una forza dissipativa di smorzamento. Il secondo principio della dinamica ci permette di scrivere l'equazione differenziale del moto dell'elettrone, da cui, ponendo

$$\mathbf{J} = -nq\mathbf{v}$$

in cui $\mathbf{J}$ è il vettore della densità di corrente, n è la concentrazione di elettroni che concorrono alla corrente, q è la carica elementare ($1,6 \times 10^{-19}$ C) e $\mathbf{v}$ il vettore velocità, si può scrivere (eq.1)

$$\frac{d\mathbf{J}}{dt} + \tau^{-1}\mathbf{J} = -\frac{nq^2}{m}\mathbf{E}.$$

Il significato del termine τ (delle dimensioni di un tempo) può essere compreso con il seguente ragionamento: se considero un'eccitazione nulla ($\mathbf{E}=\mathbf{0}$), come si verifica ad esempio subito

dopo la cessazione dell'eccitazione stessa, l'equazione diventa

$$\frac{d\mathbf{J}}{dt} + \tau^{-1}\mathbf{J} = 0$$

e descrive il decadimento della corrente transitoria nel materiale. La soluzione di quest'equazione è

$$\mathbf{J} = \mathbf{J}_0 e^{-t/\tau}$$

in cui $\mathbf{J}_0$ è il vettore della densità di corrente nell'istante in cui cessa l'eccitazione. Da questo risultato si comprende che τ è il tempo necessario affinché la corrente diminuisca fino ad una frazione pari ad e^{-1} del suo valore iniziale.

Nel caso in cui il campo elettrico sia statico (ossia non vari nel tempo), l'equazione (eq.1) diventa

$$\tau^{-1}\mathbf{J} = \frac{nq^2}{m}\mathbf{E}$$

Possiamo scrivere la conducibilità a pulsazione nulla (ovvero a frequenza nulla - in corrente continua), che è il rapporto tra i moduli della densità di corrente e del campo elettrico:

$$\sigma_0 = \frac{\mathbf{J}}{\mathbf{E}} = \frac{nq^2}{m}\tau.$$

Si ricorda che la pulsazione ω è strettamente legata alla frequenza f dalla relazione $\omega=2\pi f$.

Qualora fossimo in condizioni non stazionarie, condizione tipica dei segnali con contenuto informativo, l'analisi procede considerando eccitazioni di tipo sinusoidale, sfruttando perciò il formalismo dei vettori complessi (estensione dei fasori ai vettori, possono vedersi come vettori che hanno come coordinate dei fasori). Tale procedimento è giustificato dal fatto che, come

sappiamo, un generico segnale può considerarsi, sotto opportune condizioni, come una sovrapposizione di sinusoidi alle frequenze tipiche dello spettro del segnale. In questo caso si può scrivere la (eq.1) in questa maniera:

$$(-j\omega + \tau^{-1})\mathbf{J} = \frac{ne^2}{m}\mathbf{E} = \tau^{-1}\sigma_0\mathbf{E}$$

per cui, in questo caso

$$\mathbf{J} = \frac{\sigma_0}{1 - j\omega\tau}\mathbf{E}$$

e quindi la conducibilità è

$$\sigma(\omega) = \frac{\sigma_0}{1 - j\omega\tau}$$

E' evidente che la conducibilità ha dipendenza dalla frequenza, pertanto possiamo dire che la proprietà di un determinato materiale di condurre l'elettricità dipende anche dalla frequenza dell'eccitazione, ovvero dalle frequenze del segnale che deve attraversarlo. **Si noti che tale modello è adatto solo per materiali non isolanti, per i quali possiamo considerare elettroni di conduzione non vincolati ai nuclei degli atomi.**

Una volta note le relazioni che governano la conducibilità, possiamo focalizzare l'attenzione sulla proprietà di trasparenza. **Quando diciamo che un oggetto è trasparente?** Generalmente un oggetto si dice trasparente quando si può vedere, attraverso di esso, ciò che vi è posto dietro. Una tale definizione di trasparenza chiama in causa non solamente leggi fisiche ma anche (almeno) la fisiologia legata alla visione nell'uomo. Sfortunatamente, per trovare criteri utili al

progetto, ci soffermiamo solo allo studio delle proprietà ottiche del materiale, prescindendo dalle caratteristiche del sistema visivo umano e degli apparati che con esso collaborano. Considereremo trasparente un oggetto se, nello spettro del visibile, il campo elettromagnetico trasmesso prevale sul campo elettromagnetico riflesso dall'oggetto. In questo modo possiamo misurare la trasparenza di una sostanza attraverso misure dei coefficienti di riflessione e di trasmissione del campo, ovvero misurando l'intensità dei campi riflesso e trasmesso.

Prima di entrare nel merito della questione è necessario un passaggio preliminare. In relazione al loro comportamento elettrico, possiamo fare una distinzione tra i materiali conduttori metallici e i materiali dielettrici. I conduttori metallici si lasciano attraversare da una corrente di elettroni (conduzione ohmica o di prima specie) ma schermano, riflettendolo, il campo elettromagnetico; i dielettrici, al contrario, si lasciano permeare dal campo elettromagnetico ma non permettono una corrente elettrica al loro interno. Fortunatamente *natura non facit saltus*, perciò disponiamo anche di materiali semiconduttori, che presentano caratteristiche intermedie. E' proprio nei semiconduttori allora che cerchiamo un materiale che permetta sia una buona corrente elettrica che una buona trasmissione del campo elettromagnetico. La conducibilità dei materiali semiconduttori può essere incrementata attraverso un opportuno **drogaggio**, ovvero un processo di impiantazione di atomi che introducano nel semiconduttore degli elettroni in eccesso (oppure delle lacune, ossia mancanze di elettroni, che possono considerarsi cariche positive) che possano migliorare il meccanismo di conduzione, rendendo, al limite, il semiconduttore del tutto analogo a un metallo

(in quest'ultimo caso diciamo **degenere** il semiconduttore). In questo articolo si fa riferimento a un semiconduttore drogato "di tipo n", ovvero con eccesso di elettroni. Nei materiali conduttori, la permittività elettrica può scriversi, sulla base delle precedenti considerazioni sulla composizione sinusoidale dei segnali, come:

$$\epsilon_c(\omega) = \epsilon(\omega) + \frac{\sigma(\omega)}{j\omega}$$

che per alte frequenze può scriversi nella forma

$$\epsilon_{c\infty}(\omega) = \epsilon_{rc\infty}(\omega)\epsilon_0 = \epsilon_{\infty}(\omega) + \frac{\sigma_0}{j\omega + \omega^2\tau} \approx \epsilon_{\infty}(\omega) + \frac{\sigma_0}{\omega^2\tau}$$

(si ricorda che in un materiale in cui avvenga conduzione è definita la permittività elettrica complessa come $\epsilon_c = \epsilon + \sigma/j\omega$).

Dalla precedente formula, introducendo la grandezza ω_p , nota come pulsazione di plasma e pari a

$$\omega_p = \sqrt{\frac{\sigma_0}{\epsilon_0\tau}} = \sqrt{\frac{nq^2}{\epsilon_0m^*}}$$

in cui m^* è la massa efficace dell'elettrone e n la concentrazione di elettroni (introdotti tramite drogaggio nel caso dei semiconduttori), posso ottenere la scrittura

$$\epsilon_{rc\infty}(\omega) = \epsilon_{\infty}(\omega) + \frac{\omega_p^2}{\omega^2 + j\frac{\omega}{\tau}}$$

Facciamo un po' di chiarezza sui precedenti risultati: la permittività elettrica è un indice di quanto il materiale tenda ad opporsi ad un campo elettrico esterno, indebolendolo; si nota che essa è legata alla pulsazione di plasma ω_p - è

possibile, al solito, definire una frequenza di plasma $f_p = 2\pi\omega_p$. Per frequenze superiori alla sua frequenza di plasma, un materiale ha comportamento dielettrico, al di sotto di tale frequenza, al contrario, il materiale riflette il campo elettromagnetico.

Dal modello di Drude si ricava anche che

$$\sigma_0 = nq\mu$$

dove μ indica la mobilità elettronica, una grandezza legata alla velocità degli elettroni e di cui parleremo a breve.

Dalle due precedenti relazioni vediamo che la conducibilità e

la permittività dipendono dalla concentrazione degli elettroni nel materiale, queste considerazioni hanno portato in un primo momento ad attuare drogaggi più spinti per aumentare la conducibilità del materiale: al crescere del numero di portatori, infatti, cresce la conducibilità; sfortunatamente, per effetto collaterale, si riduce la mobilità elettronica, fatto che incide negativamente sulla conducibilità. In seconda istanza si è visto sperimentalmente che aumentando il drogaggio non si ha un pari aumento di portatori e in più si riduce la trasparenza del film.

La tendenza dunque è: aumentando inizialmente il numero di elettroni si migliora la conducibilità ma si peggiora la trasparenza del materiale; eccedendo con l'introduzione di elettroni si peggiora anche la mobilità elettronica, con conseguenze negative sulla conducibilità. Il compromesso è da cercare nel drogaggio che massimizzi la conducibilità senza alterare troppo la trasparenza e la mobilità dei portatori.

Per far combaciare una buona conducibilità alla trasparenza del materiale, deve accadere che

la lunghezza d'onda corrispondente alla frequenza di plasma sia maggiore della minima lunghezza d'onda visibile dall'occhio umano. Conservando un minimo di margine su quest'ultima, si può scrivere

$$\frac{2\pi}{\omega_p} > \frac{\lambda'_{vis}}{c}$$

da cui ricaviamo una condizione sulla concentrazione di portatori:

$$n < \frac{4\pi^2 \epsilon_\infty m^*}{\mu_0 q^2 \lambda'_{vis}{}^2}$$

Questa condizione si configura come un **vincolo da rispettare sul drogaggio del materiale**.

LA MOBILITÀ DEI PORTATORI

Si è visto che **incrementando il numero di portatori per unità di volume, non è possibile ottenere un'elevata conducibilità senza ridurre la trasparenza del materiale**. Nei casi applicativi di interesse, i materiali vengono depositati tramite un processo denominato sputtering, in sottili strati (film) conduttivi; lo spessore ridotto del film favorisce la sua trasparenza.

Ricordiamo che $\sigma_0 = nq\mu$. Se non è possibile aumentare la concentrazione n di elettroni, **la strategia da seguire per migliorare le proprietà conduttive è cercare di ottenere alta mobilità dei portatori**. Per definizione si ha $\mu = v/E$, dal modello di Drude ricaviamo invece che $\mu = q\tau/m$. La mobilità elettronica è molto influenzata dal disordine dovuto alla particolare struttura del materiale e dalla modificazione della struttura risultante dal drogaggio (che incidono sul fattore τ).

Il moto degli elettroni all'interno di una struttura disordinata comporta un forte scattering elettronico, ovvero un gran numero di urti tra elettroni e reticolo del materiale ostacola il moto libero dei portatori stessi. Esistono diverse cause di scattering. In un materiale policristallino (ovvero formato da tanti piccoli cristalli ordinati), laddove il cammino libero medio sia paragonabile alla dimensione dei grani cristallini, si ha scattering presso i bordi dei medesimi grani; d'altronde aumentando la superficie dei grani, la conducibilità di un materiale policristallino diminuisce. Si può trascurare lo scattering superficiale (dovuto alle imperfezioni presenti su tutte le interfacce tra materiali differenti) se il cammino libero medio non è paragonabile allo spessore del film. Fin tanto che si possa considerare il materiale un semiconduttore degenere, ricoprono un ruolo di primaria importanza le impurità del cristallo come centri di scattering. Bisogna ricordare che le proprietà elettriche, come la mobilità, hanno inoltre dipendenza dalla temperatura. La mobilità risultante da cause indipendenti di scattering può trovarsi ponendo:

$$\frac{1}{\mu_{tot}} = \frac{1}{\mu_n} + \frac{1}{\mu_i}$$

(nella fattispecie si stanno considerando come cause di scattering impurità neutre ed impurità ionizzate all'interno del materiale).

Tra i materiali che possiedono le caratteristiche necessarie di conducibilità e trasparenza rientrano i **Transparent Conducting Oxide (TCO)**, ossidi trasparenti conduttivi, molto usati in optoelettronica. Nel caso delle antenne, tra i TCO si è distinto l'**ossido di indio-stagno** (in inglese **Indium-Tin Oxide (ITO)**), il materiale di gran lunga più utilizzato nella sua categoria.

L'INDIUM-TIN OXIDE

L'ITO è una lega che si ottiene dall'**ossido di indio (In_2O_3)**. Sappiamo che i materiali semiconduttori a livello microscopico sono caratterizzati da due bande di energia relative agli stati in cui gli elettroni si possono trovare. La banda a maggiore energia è detta "di conduzione", la banda a minore energia è detta "di valenza". Tra le due bande vi è una banda proibita, la cui ampiezza è detta *band-gap*. L'ossido di indio è un materiale semiconduttore ad elevato *band-gap* (la differenza di energia tra banda di conduzione e banda di valenza è di circa 3eV, più del doppio del valore che si trova nel silicio monocristallino usato per i transistor in microelettronica). Le proprietà elettriche di un film di ossido dipendono in gran parte dalla presenza di atomi di impurità (che introducono livelli energetici disponibili all'interno della banda proibita). Il reticolo dell'ossido di indio può presentare, per via dei processi di deposizione, molte vacanze di ioni O^{2-} , questo tipo di imperfezione introduce una ulteriore banda di livelli permessi a ridosso della banda di conduzione. In queste condizioni si dice che il semiconduttore è degenero e permette una conduzione di tipo ohmico. Per migliorare le proprietà conduttive si procede al drogaggio del materiale mediante introduzione di atomi di stagno, che nel reticolo vanno a sostituire in alcuni siti gli atomi di indio. Lo stagno ha un elettrone in più rispetto all'indio, questo significa che per ogni atomo di stagno verrà introdotto anche un elettrone in eccesso. In queste condizioni diciamo che il materiale è drogato di tipo n, ovvero che c'è al suo interno un eccesso di elettroni. Ne deriva una soluzione solida in cui il 90% in peso è di ossido di sta-

gno, il restante 10% invece è di ossido di indio. Il drogaggio comporta, oltre all'introduzione di elettroni, la variazione della banda proibita: in questo caso passiamo ad un *band-gap* di circa 3,8eV. Sappiamo che gli elettroni, in un materiale semiconduttore, possono passare dalla banda di valenza alla banda di conduzione, per dar vita a una corrente, se vengono investiti da un fotone che porti energia pari all'ampiezza della banda proibita. Nel caso dell'ITO, l'ampliamento della banda proibita ci allontana dal pericolo di assorbimento delle frequenze della luce visibile. Quando il semiconduttore è degenero (con l'ITO succede per drogaggi a concentrazione dell'ordine di 10^{19}cm^{-3}) il suo comportamento può spiegarsi bene tramite il modello di Drude. Per quanto riguarda la struttura, dalle deposizioni mediante sputtering si ottiene un reticolo non ordinato ma policristallino, in alcuni casi lo si può considerare addirittura amorfo.

Note le caratteristiche dell'ITO e i vincoli da rispettare per ottenere buoni risultati nelle proprietà di interesse, possiamo proseguire il nostro percorso addentrandoci nel contesto delle applicazioni e in particolar modo vediamo quali sono i fattori da tenere in conto nel progetto di un'antenna da realizzare con ITO.

Le antenne che prenderemo in considerazione dovranno essere realizzabili mediante un film di materiale, perciò ci concentreremo sulle antenne stampate. Tra queste è di particolare interesse il caso delle antenne a *patch* (delle quali già si sono realizzati numerosi esemplari in ITO in ambito sperimentale), [che trovano applicazione nelle comunicazioni wireless](#).

Proseguiamo interessandoci all'individuazione dei parametri utili al progetto.

ASSORBIMENTO ELETTROMAGNETICO DEL FILM

Al fine di ottenere il parametro di drogaggio n guardiamo al modello di assorbimento del materiale.

La propagazione di un campo all'interno di un materiale ha andamento proporzionale a $e^{-\alpha z}$, con α costante di attenuazione e z coordinata spaziale nella direzione

di propagazione. Se chiamiamo spessore della pelle δ_p la distanza percorsa dall'onda nel materiale finché non è ridotta ad e^{-1} del suo valore di superficie,

possiamo scrivere

$$e^{-\alpha z} = e^{-1}$$

$$-\alpha z = -1$$

$$\delta_p = z$$

allora, ponendo la definizione del numero d'onda $k = \alpha + j\beta$, scrivo

$$\delta_p = \frac{1}{\alpha} = -1/\Im\{k\}$$

Nello spettro in cui il materiale è trasparente (ovvero per frequenze superiori alla frequenza di plasma ma non tali che i fotoni possano essere assorbiti dagli elettroni per il passaggio dalla banda di conduzione a quella di valenza) si può trovare un'espressione più comoda:

$$\delta_p \approx \frac{2m^*}{Z_\infty q^2} \frac{\omega^2 \tau}{n}$$

con

$$Z_\infty = \frac{377}{\sqrt{\epsilon_\infty}} \Omega.$$

Sperimentalmente si può vedere che il coefficiente di trasmissione per la luce ha un'espressione del tipo:

$$T(t) \approx e^{\frac{-2t}{\delta_p}}$$

Se pongo la permittività elettrica complessa ad alta frequenza:

$$\epsilon_{c\infty}(\omega) = \epsilon_1 - j\epsilon_2$$

allora posso scrivere

$$\epsilon_1 = \epsilon_\infty - \frac{\omega_p^2 \tau^2}{1 + \omega^2 \tau^2} \approx \epsilon_\infty$$

$$\epsilon_2 = \frac{1}{\omega} \frac{\omega_p^2 \tau}{1 + \omega^2 \tau^2} \approx \frac{\omega_p^2 \tau}{\omega^3} \ll \epsilon_1 \approx \epsilon_\infty.$$

In particolare, dalla letteratura si possono trovare come valori di riferimento i seguenti:

$$\epsilon_\infty \approx 4,0$$

$$m^* \approx 0,35 m_e$$

$$\tau \approx 3,3 \cdot 10^{-15} \text{ s.}$$

dove m_e indica la massa a riposo dell'elettrone. Per i film di buona qualità si riescono ad ottenere $\mu \approx 50 \text{ cm}^2 \text{ V}^{-1} \text{ s}^{-1}$ e $\tau \approx 10^{-14} \text{ s}$. Da queste considerazioni si deduce che, mantenendo un margine di $1 \mu\text{m}$ sulla lunghezza d'onda, la massima concentrazione impiantabile di portatori è $n_{\text{max}} = 1,5 \cdot 10^{21} \text{ cm}^{-3}$.

LA RESISTENZA SUPERFICIALE

La resistenza all'interno del film determinerà l'entità delle dissipazioni di energia per effetto joule, influenzando negativamente sull'efficienza dell'antenna.

Se consideriamo un film con conducibilità σ in presenza di un campo elettrico variabile con lunghezza d'onda λ (ricordiamo che la lunghezza d'onda è legata alla frequenza dalla relazione $\lambda = c/f$, dove c è la velocità della luce), il materiale sarà soggetto ad una densità di corrente

$$\mathbf{J} = \sigma \mathbf{E}'$$

dove $\mathbf{E}'$ indica il campo elettrico tangenziale alla superficie all'interno del materiale. Se lo spessore del film t è molto inferiore alla lunghezza d'onda, possiamo pensare che la corrente sia tutta superficiale:

$$\mathbf{J}_s = t\mathbf{J}$$

per cui

$$\mathbf{E}' = (\sigma t)^{-1} \mathbf{J}_s.$$

Poiché il campo deve essere continuo alla superficie di separazione con l'ambiente, vale:

$$\mathbf{E}_{\text{tan}} = (\sigma t)^{-1} \mathbf{J}_s.$$

Definiamo allora la resistenza superficiale R_s come il rapporto tra i moduli del campo elettrico tangenziale alla superficie e della densità di corrente superficiale:

$$R_s = (\sigma t)^{-1}.$$

I produttori, tipicamente, specificano le proprietà elettriche dei film sottili proprio mediante la

resistenza superficiale e lo spessore del film, resta allora univocamente determinata la conducibilità, parametro che utilizziamo in **fase di progettazione** attraverso l'utilizzo di software simulatori elettromagnetici.

L'EFFICIENZA D'ANTENNA

Un'antenna è disegnata per ospitare delle correnti elettriche e per reagire ai campi elettromagnetici, se il materiale oppone resistenza, si hanno perdite dovute a **dissipazioni ohmiche**. Un'altra forma di dissipazione è data dalle **perdite dielettriche** ed è dovuta al fatto che una parte del campo viene "sprecata" per polarizzare il materiale (si verificano rotazioni e deformazioni a livello microscopico). Abbiamo già visto che la resistenza superficiale è legata alla conducibilità del materiale e allo spessore del film, vediamo ora come la resistenza incide sull'efficienza dell'antenna.

Considerando l'antenna in trasmissione, l'efficienza è direttamente dipendente da più fattori:

1. grado di adattamento tra generatore ed antenna;
2. perdite dielettriche ed ohmiche all'interno dell'antenna;
3. polarizzazione rispetto all'antenna ricevente.

Ad ognuno di questi fattori, presi separatamente, può essere associato un fattore di efficienza. L'efficienza totale dell'antenna è pari al prodotto dei fattori di efficienza.

Tra questi fattori, solo quello dovuto alle perdite è direttamente vincolato dalla struttura fisica dell'antenna. Si supponga che l'inefficienza sia dovuta solamente a perdite ohmiche e dielettriche all'interno di un conduttore. L'efficienza dell'antenna, nota in tal caso come efficienza di radiazione, è definita come il rapporto tra la

potenza trasmessa, ovvero irradiata, W_T e la potenza disponibile del generatore

$$\eta = \frac{W_T}{W_g}$$

dove generalmente

$$W_g = \frac{1}{2} \frac{V_g^2}{R_g}$$

con R_g ad indicare la resistenza interna del generatore e V_g ad indicare il valore massimo della tensione imposta dal generatore. Si associa inoltre la potenza irradiata dall'antenna ad una resistenza di radiazione R_r , mentre la potenza dissipata è associata ad una resistenza di perdita R_L . Le resistenze sono legate alle rispettive potenze dalle relazioni:

$$W_T = \frac{1}{2} |I_g|^2 R_r$$

$$W_L = \frac{1}{2} |I_g|^2 R_L$$

dove $|I_g|$ indica il valore massimo della corrente che attraversa l'antenna. Si ricorda che, nel caso considerato, in cui l'inefficienza è dovuta solo a perdite ohmiche e dielettriche, deve essere verificata l'uguaglianza

$$W_g = W_T + W_L.$$

Si può allora riscrivere l'efficienza nella forma:

$$\eta = \frac{W_T}{W_T + W_L}$$

dalla quale, sostituendo le espressioni trovate in luogo delle potenze irradiata e dissipata, si ottiene l'utile relazione:

$$\eta = \frac{R_r}{R_r + R_L}.$$

Poiché R_L è indice delle perdite dissipative, il materiale di cui è costituita l'antenna è un fattore molto importante ai fini di una buona efficienza di radiazione. L'efficienza di radiazione, a sua volta, è un parametro di importanza primaria poiché indice di prestazioni. Quando si vuole descrivere un'antenna in termini di prestazioni, infatti, si fa riferimento al **guadagno**, che è il prodotto tra l'efficienza d'antenna e la direttività dell'antenna e descrive quanta della potenza in ingresso all'antenna viene irradiata nella regione di massima irradiazione (e quindi dipende da quanta potenza viene invece dissipata all'interno dell'antenna).

CONCLUSIONI

Nella pratica attualmente si riescono ad ottenere antenne per applicazioni wireless (attorno ai 5.8 GHz ma anche fino a 10GHz) con guadagni di quasi 0dB e trasparenze di oltre il 70% (ovvero meno del 30% della luce incidente viene riflessa).

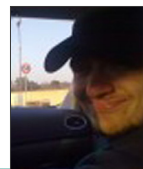
Come molti intenderanno, prestazioni di questa entità sono piuttosto deludenti se paragonate a quelle delle antenne realizzate in alluminio o in rame; in ambito applicativo, tuttavia, qualche realizzazione sembra aver trovato la sua strada. Un esempio è quello delle antenne trasparenti stampate su specchi, in grado di comunicare con chip **RFID**, per applicazioni in negozi di abbigliamento (un cliente potrebbe disporre in tempo reale delle informazioni sul capo che

sta provando in camerino, senza accorgersi delle antenne stampate sullo specchio); in questo specifico caso, le antenne comunicano a 2.5GHz e, per via dell'uso di un array di antenne e di una scelta intelligente della geometria, si riesce ad ottenere un guadagno ben superiore ai "miseri" 0dB ottenuti nel migliore degli altri casi tipici (con singola antenna).

Il futuro delle antenne trasparenti si gioca ora sul campo dei materiali, sebbene l'ITO sia stato il più gettonato finora, altre leghe come l'ossido di indio-fluoro (**Fluorine-Tin Oxide (FTO)**) sono oggetto di ricerca. Un altro tipo di antenne trasparenti è quello delle antenne "a maglie" (*meshed*), che con una filosofia leggermente diversa riesce ad ottenere un risultato analogo ma più prestante - infatti il numero di applicazioni in questo caso è nettamente superiore. Allo stato attuale dell'arte, si lascia intendere che molte delle speranze sono riposte nei futuri sviluppi delle scienze dei materiali e dei processi della microelettronica, fino ad allora lasciamo lavorare gli ingegneri affinché trovino le soluzioni migliori con i mezzi odierni.

L'autore è a disposizione nei commenti per eventuali approfondimenti sul tema dell'Articolo.
Di seguito il link per accedere direttamente all'articolo sul Blog e partecipare alla discussione:
<http://it.emcelettronica.com/antenne-trasparenti-la-tecnologia-del-mimetismo>

LTE: scopriamo la telefonia mobile 4G



di gransasso
17 Settembre 2015



LTE (Long Term Evolution) è l'ultima evoluzione per ciò che riguarda la rete radio-mobile a diffusione mondiale. Lanciata dieci anni dopo l'UMTS e venti dopo il GSM, LTE si caratterizza per la grande velocità di trasmissione ed il ridotto tempo di latenza rispetto alle precedenti tecnologie radiomobili. In questo articolo analizzeremo le differenze tecnologiche di LTE rispetto al sistema UMTS, non prima però di aver introdotto l'architettura del sistema cellulare. Inoltre verrà affrontata l'evoluzione storica del sistema radiomobile in Italia in modo da dare anche ai meno esperti la possibilità di comprendere le differenze e le migliorie apportate nel corso degli anni.

STORIA DEL SISTEMA RADIOMOBILE IN ITALIA

La storia del sistema radiomobile in Italia è rappresentativa dell'evoluzione tecnologica di questo tipo di sistemi a livello mondiale. In particolare prima dell'avvento dei sistemi universali, dalla terza generazione in poi, Stati Uniti, Europa e paesi Asiatici si sono mossi su

percorsi paralleli ma distinti ognuno con i propri standard e le proprie differenze.

In Italia il primo sistema radiomobile di massa fu il TACS (*Total Access Communication System*), la cosiddetta generazione 1. Lanciato nel 1990 si basa su una comunicazione di tipo analogico a commutazione di circuito con modalità di accesso a divisione di frequenza (FDMA) e

modulazione FM ed operante nella gamma dei 450 MHz. Nel 1993 viene introdotto in Italia l'E-TACS (*Enhanced Total Access Communication System*) operante nella gamma dei 900 MHz.

Nel 1995 avviene il lancio della seconda generazione ossia il sistema GSM (*Global System for Mobile Communications*) e per la prima volta si sente parlare di SMS. GSM sfrutta la comunicazione digitale a commutazione di circuito, con modalità di accesso a divisione di tempo (TDMA), con modulazione del tipo *Shift Keying* ed operante nella gamma dei 900 MHz. Con l'aumento del traffico e la naturale scarsità di risorse, nel 1998 si comincia a trasmettere anche nella gamma dei 1800 MHz giungendo al cosiddetto *dual band*. Nel 2001 viene abilitata la commutazione di pacchetto grazie alla generazione 2.5 denominata GPRS (*General Packet Radio Service*) che consente l'accesso ad Internet, e nel 2003 la sua evoluzione EDGE (*Enhanced Data rates for GSM Evolution*) ne aumenta la velocità di trasferimento dati.

Nello stesso anno si comincia a diffondere la terza generazione ossia UMTS (*Universal Mobile Telecommunications System*). Basato su comunicazione digitale a commutazione di pacchetto, con modalità di accesso a divisione di codice (CDMA), modulazione QPSK ed operante nella gamma dei 2100 MHz. Grazie all'aumento della velocità di trasmissione ed alla migliore efficienza in banda rispetto agli standard precedenti, videochiamate e navigazione in internet su mo-

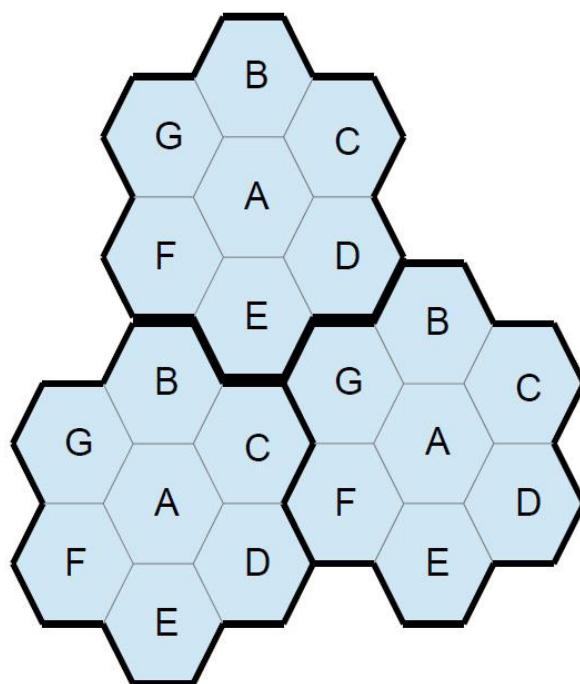
bile iniziano ad essere all'ordine del giorno. Nel 2006 la prima evoluzione HSPDA (*High Speed Downlink Packet Access*) e nel 2009 la seconda evoluzione HSPDA+, permettono di aumentare la velocità di trasmissione di DL (*downlink*) e UL (*uplink*).

Si giunge infine al 2012 con l'avvento della cosiddetta quarta generazione LTE (*Long Term Evolution*). Di LTE si parlerà ampiamente nel paragrafo intitolato "Comparazione tra UMTS e LTE". Di seguito viene invece utilizzata una tabella per riassumere quanto detto sull'evoluzione del sistema radiomobile cellulare in Italia.

Sistema	Anno rilascio	Commutazione	modalità di accesso	modulazione	gamma	banda
TACS	1990	analogica circuito	FDMA	FM	450 MHz	25 KHz
E-TACS	1993	analogica circuito	FDMA	FM	900 MHz	25 KHz
GSM	1995	digitale circuito	TDMA	SK	900 MHz	200 KHz
GPRS	2001	digitale pacchetto	TDMA	SK	1800 MHz	200 KHz
EDGE	2003	digitale pacchetto	TDMA	SK	1800 MHz	200 KHz
UMTS	2003	digitale pacchetto	WCDMA - FDD	QPSK	2100 MHz	5 MHz
HSPDA	2006	digitale pacchetto	WCDMA - FDD	QPSK	2100 MHz	5 MHz
HSPDA+	2009	digitale pacchetto	WCDMA - FDD	QPSK	2100 MHz	5 MHz
LTE	2012	digitale pacchetto	OFDM	QAM	4 bande	5 MHz

ARCHITETTURA SISTEMA CELLULARE

Un sistema di telefonia cellulare consente la comunicazione wireless tra gli utenti presenti all'interno delle aree coperte dal sistema. Il territorio coperto dalla rete viene suddiviso in *celle* di dimensioni variabili. Ogni cella è "servita" da una *stazione radio base (BS)* cioè una torre dove sono posizionate le antenne per trasmettere e ricevere i segnali. Lo spettro totale di frequenze disponibili viene suddiviso in gruppi, ed a celle adiacenti vengono assegnate gruppi di canali differenti in modo da minimizzare le interferenze e permettere il *riuso delle frequenze* su tutto il territorio. Nella figura in basso viene proposto un esempio del concetto di *riuso delle frequenze*. Celle con la stessa lettera usano lo stesso gruppo di frequenze. Le celle raggruppate con un bordo nero spesso vengono definite *cluster* ed è ciò che viene replicato sul territorio da coprire.



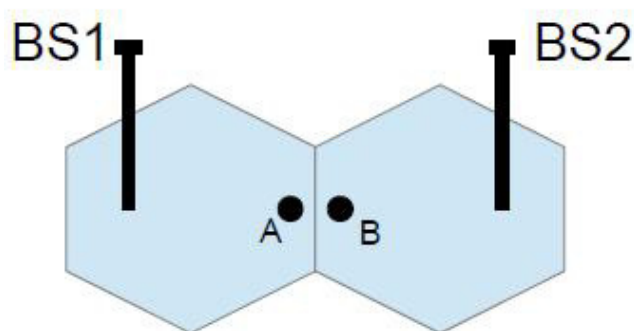
La forma esagonale di ogni singola cella, come visibile dalla figura sopra, è puramente concettuale, ma è universalmente adottata poichè permette una facile analisi di un sistema cellulare. Normalmente per antenne omnidirezionali la stazioni radio base sono situate al centro della cella mentre nei casi di antenne settoriali le torri

vengono posizionate agli angoli della cella.

In precedenza si è accennato al fatto che le celle sono di dimensioni variabili, questo è funzione della densità di utenza da servire nell'area geografica. Senza passare attraverso formule, ma semplicemente con un ragionamento euristico: se il traffico massimo smaltibile da una cella è X , allora più la densità di popolazione aumenta più piccola sarà la cella, questo per favorire il riuso delle frequenze e così servire un maggior numero di utenti. Da ciò si può capire che un'area urbana sarà coperta da molte *micro celle*, mentre un'area rurale delle stesse dimensioni la si potrà ricoprire anche con una singola *macro cella*. Le celle vengono così distinte in base alla loro dimensione e disponendole in ordine decrescente si avranno: macrocelle, microcelle, picocelle e femtocelle.

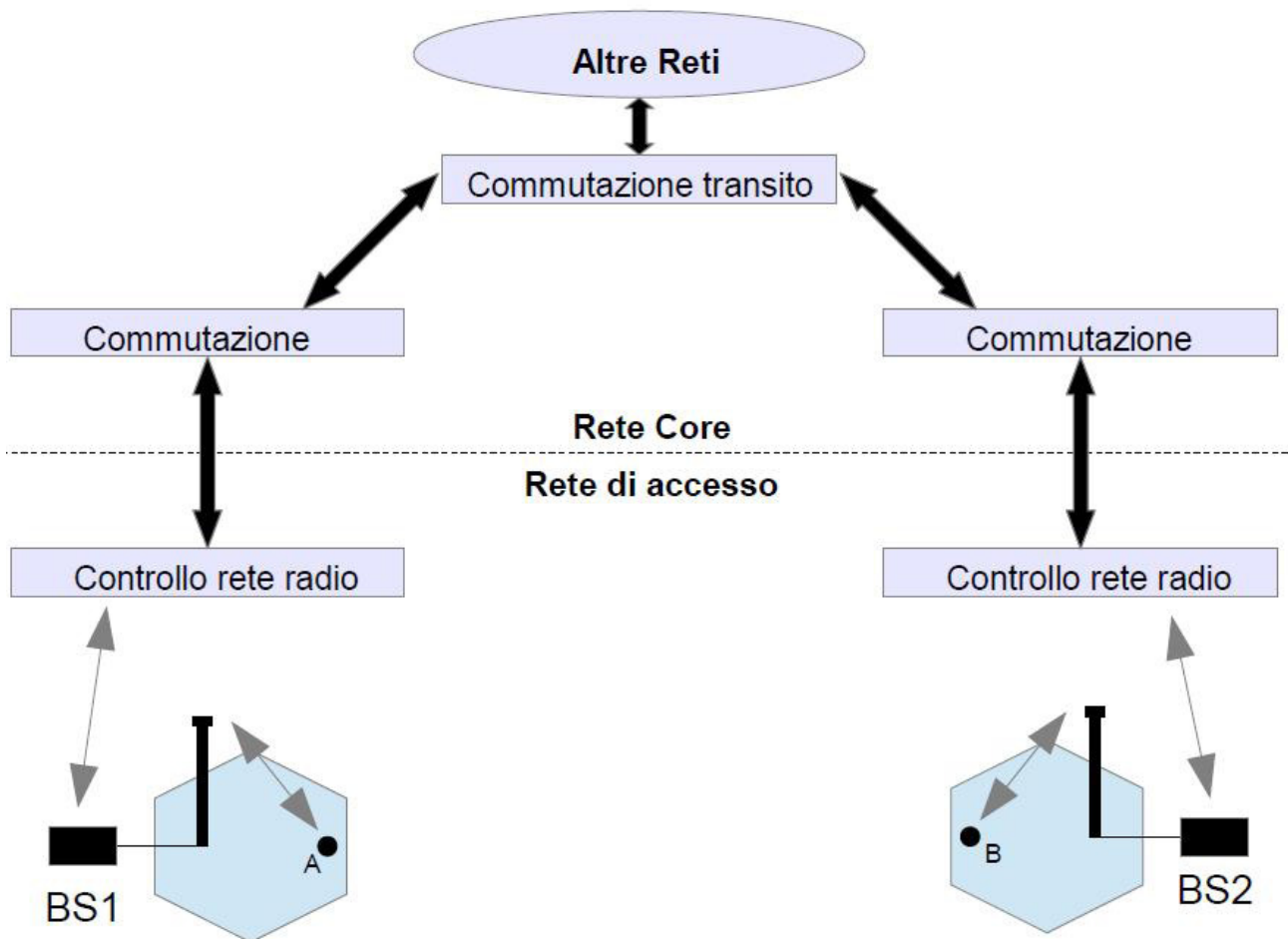
La telefonia mobile è tale poichè permette di muoversi usufruendo dei servizi offerti. Poichè celle adiacenti utilizzano canali differenti, se durante una telefonata si attraversasse il confine tra due celle la comunicazione dovrebbe cadere. Questo non accade grazie ad un meccanismo chiamato *handover*. La figura in basso mostra un esempio di *handover*: nel passaggio dal punto A al punto B la comunicazione deve essere allocata sui canali della stazione radio base BS2, prima che la potenza del segnale ricevuto da BS1 scenda al di sotto di una soglia minima per mantenere la chiamata. Allontanandosi dal punto A la potenza del segnale cala, quindi il protocollo identifica una nuova BS con potenza del segnale in crescita ed inizia il meccanismo di passaggio della chiamata da un canale all'altro. Il tutto è naturalmente calibrato per tenere in conto molti fattori tecnici che però non menzio-

neremo in questo articolo.



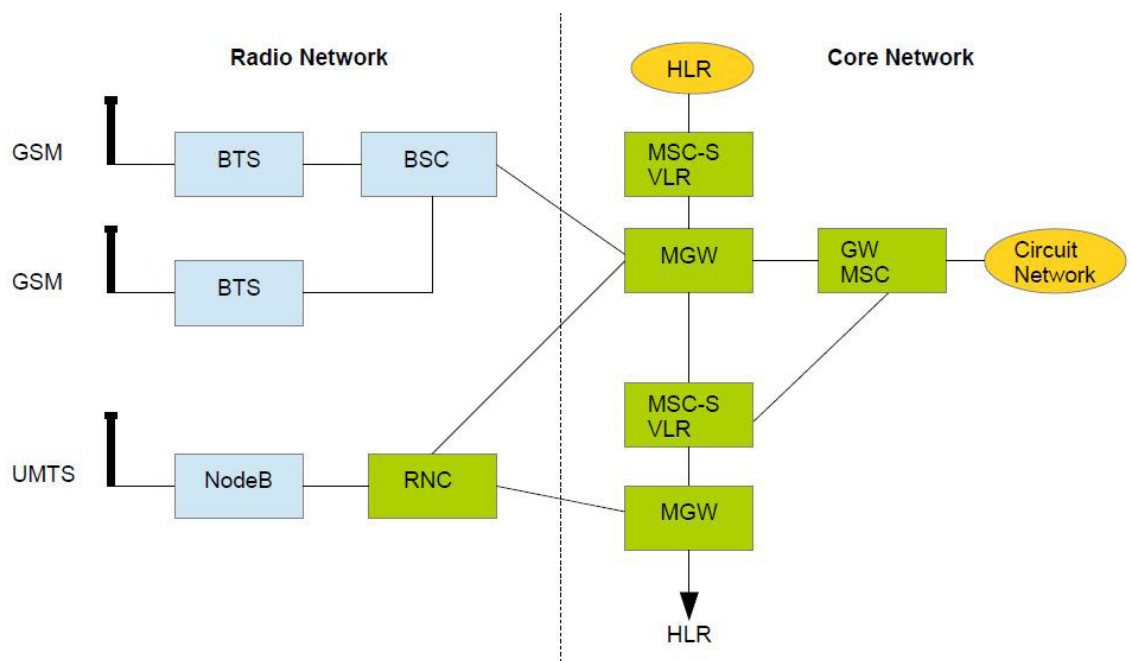
A questo punto, dopo questi principi di base, occorre salire di livello dal punto di vista architetturale. Il diagramma a seguire mostra l'architettura generale di un sistema radiomobile cellulare. Per prima cosa si può evidenziare la distinzione tra *rete di accesso* e *rete core*. La prima è costituita dall'interfaccia radio tra terminali mobili e stazioni radio ed è caratterizzata da meccanismi di accesso multiplo. La seconda è costituita dalla parte cablata della rete che realizza e gestisce il collegamento fisico con il servizio richiesto dall'utente. Tali servizi sono quelli tipici, come ad esempio: comunicare con un utente di un'altra cella, con un utente di telefonia fissa, o il collegamento con altre reti come quella *internet*.

Il *controllo rete radio* gestisce l'assegnazione iniziale delle risorse (cioè il canale per la comunicazione), la nuova assegnazione per l'*handover* e il rilascio finale del canale usato. Il blocco *commutazione* si occupa di instradare la chiamata da utente chiamante a utente ricevente mobile o della telefonia fissa. Mentre il blocco di *transito* realizza la connessione all'interno della stessa rete dell'operatore o verso altre reti fisse o mobili.



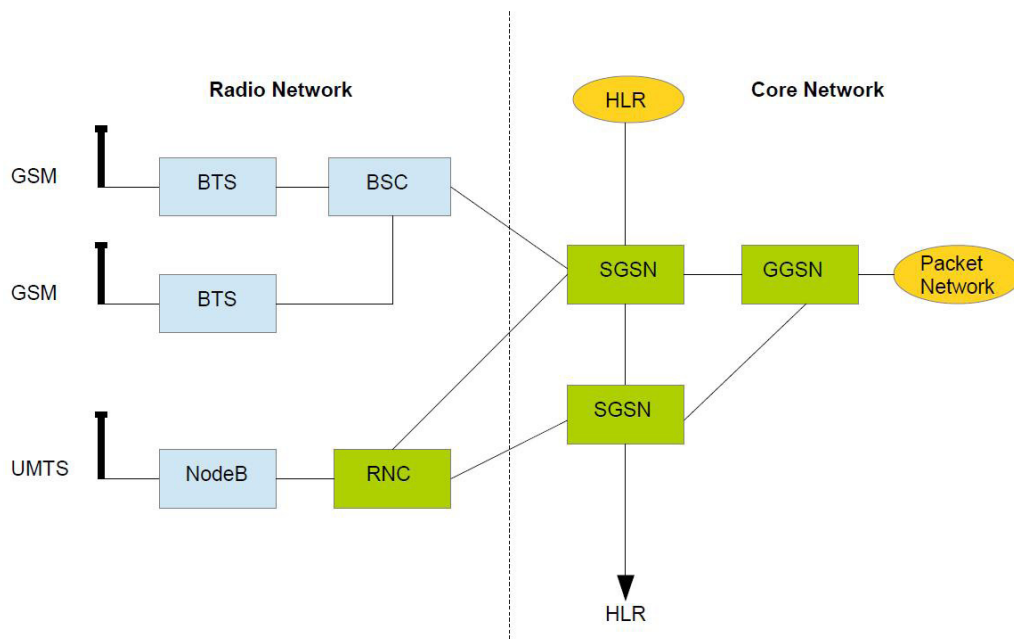
La *commutazione* può avvenire per mezzo di due tecniche: *commutazione di circuito CS* (*Circuit Switching*) in cui le risorse per la comunicazione vengono impegnate per tutta la durata della connessione, oppure *commutazione a pacchetto PS* (*Packet switching*) in cui la comunicazione si suddivide in blocchi di dati che impegnano le risorse solo per il tempo di transito di ogni singolo blocco. La rete GSM e quella UMTS, salvo le interfacce radio, sono mol-

to simili dal punto di vista architettonico, quindi è possibile utilizzare entrambe per descrivere un'architettura basata sulla commutazione di circuito. Nella figura sotto viene raffigurato lo schema a blocchi delle reti core e radio nel caso di commutazione di circuito.



La rete di accesso radio (Radio Network) è costituita da:

- Basic Transceiver Station (BTS) nel caso GSM e NodeB (NB) nel caso UMTS, che rappresentano i nodi di accesso.
- Base Station Controller (BSC) nel caso GSM e Radio Network Controller (RNC) nel caso UMTS interconnettono i nodi di accesso con la rete network. Inoltre gestiscono le risorse e l'handover.



La rete core (Core Network) è costituita da:

- Mobile Switching Center Server (MCS-S) che si occupa della commutazione di circuito e altre funzioni collaterali.
- HLR e VLR sono due registri. Il primo è un database centralizzato dove sono contenuti i dati dell'utente come autenticazione e posizione in termini di VLR. Il secondo è associato ad ogni MSC e contiene i dati degli utenti in visita. Quando viene effettuata una chiamata viene interrogato HLR per sapere il VLR dell'utente chiamato, quindi l'MSC relativo dirama una chiamata su tutte le celle da lui controllate, ed infine alla risposta dell'utente instrada la chiamata alla cella corrispondente.

Nella figura sotto, viene invece raffigurato lo schema a blocchi delle reti core e radio nel caso di commutazione di pacchetto.

La parte di accesso radio è simile a quella già vista nella commutazione di circuito. La *core network* invece si distingue per:

- SGSN (Serving GPRS Support Node) gestisce l'autenticazione e la mobilità d'utente, sia in rete di appartenenza, sia in rete visitata nel caso di roaming.
- GGSN (Gateway GPRS Support Node) consente la connessione IP per l'accesso ai servizi erogati dall'operatore di appartenenza dell'utente, dalle Corporate e da Internet.

COMPARAZIONE TRA UMTS E LTE

LTE viene definita come la quarta generazione dei sistemi radiomobili cellulari, anche se questa dicitura è di natura puramente commerciale, poichè dal punto di vista tecnico bisognerebbe parlare di pre-4G. La rete LTE si pone infatti a metà strada tra gli standard 3G (UMTS) e quelli 4G (LTE advanced) che hanno l'obiettivo di raggiungere velocità di trasferimento di 1Gb/s.

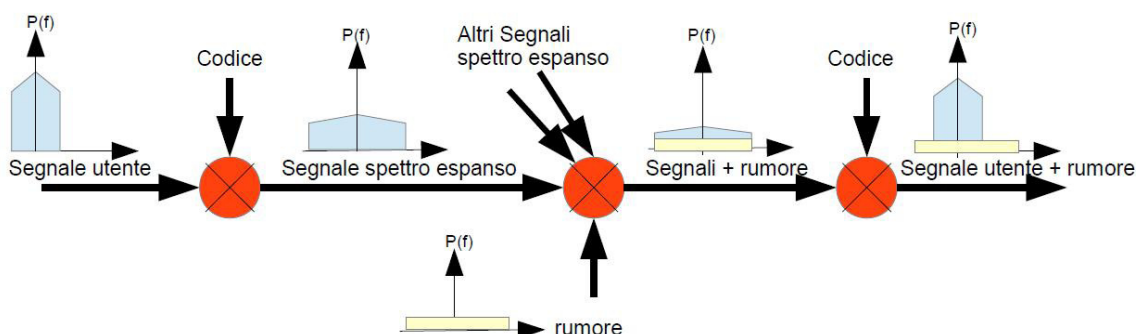
LTE può operare in 4 gamme di frequenze: 800 MHz derivanti dalla liberazione dei vecchi canali televisivi analogici, 900 MHz, 1800 MHz e 2600

MHz.

A differenza dell'UMTS che utilizza l'accesso a divisione di codice (CDMA), LTE utilizza in *downlink* (dalla stazione radio base all'utente) la modulazione OFDM (*Orthogonal Frequency-Division Multiplexing*) mentre in *uplink* (dall'utente alla stazione radio base) l'accesso multiplo a divisione di frequenza a portante singola SC-FDMA (*Single Carrier Frequency Division Multiple Access*). Questo consente una maggiore efficienza in banda (cioè maggiore bit-rate per singolo Hz), minore latenza (cioè ritardo tra invio pacchetto e ritorno) e maggiori velocità di trasferimento in downlink e in uplink. Vediamo tecnicamente cosa tutto ciò vuol dire. Nel CDMA:

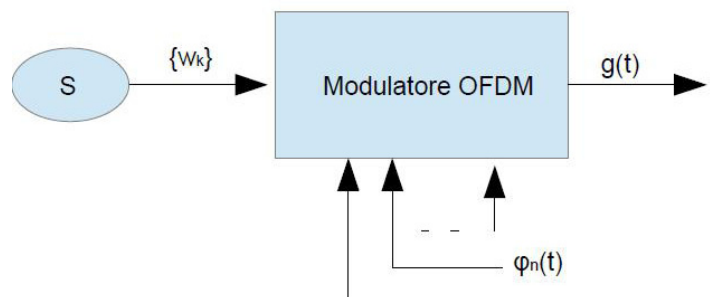
- ad ogni utente viene assegnato un codice.
- il codice "modula" il segnale dell'utente.
- i codici assegnati tra i vari utenti sono fra loro ortogonali, nel senso che il prodotto scalare tra due codice-vettori differenti è pari a zero. Quindi anche i segnali sono tra loro ortogonali.
- i segnali così modulati viaggiano tutti sullo stesso canale, cioè alla stessa frequenza.
- in ricezione il singolo utente conoscendo il suo codice può recuperare il suo segnale moltiplicandolo di nuovo per il codice. Infatti grazie all'ortogonalità, il prodotto tra segnali differenti farà zero.

Lo schema in basso riassume sinteticamente quanto detto sopra.



La modulazione del segnale utente attraverso il codice provoca il fenomeno dello *Spread spectrum*, cioè lo spettro del segnale viene spalmato su un intervallo di frequenze maggiore con conseguente appiattimento (chiaramente l'energia deve restare uguale). I segnali vanno a porsi quasi al di sotto della potenza di rumore, ed infatti questa tecnica era all'inizio utilizzata dai militari per le trasmissioni cifrate. Intuitivamente il segnale necessita di una banda superiore alla sua a causa dello *Spread Spectrum* e questo ci porta a pensare che l'efficienza in banda possa non essere delle migliori.

Nell'OFDM invece, supponiamo che il segnale sia stato modulato con una modulazione QAM avremo quindi una successione di simboli w_k indipendenti e identicamente distribuiti. Inviando i simboli, nel modulatore OFDM, essi vengono associati ognuno ad una portante $\phi_n(t)$ come nella figura sotto.



Si ottiene così un segnale $g(t)$ che è la sommatoria dei prodotti tra simboli e portanti associate:

$$g(t) = A_c \sum_{n=0}^{N-1} w_n \phi_n(t) \quad 0 < t < T$$

T rappresenta la durata del segnale modulato ed è pari al prodotto di NT_s dove N è il numero delle portanti e T_s è la durata del singolo simbolo. Le portanti sono ortogonali tra di loro con una distanza in frequenza tra loro pari a $1/T$.

$$\phi_n(t) = e^{j2\pi f_n t} \quad \text{dove} \quad f_n(t) = \frac{1}{T} \left(n - \frac{N-1}{2} \right) \quad n=0, \dots, N-1$$

Campionando il segnale $g(t)$ con passo di campionamento T_s (cioè la durata del simbolo) si ottiene attraverso una serie di passaggi matematici:

$$g'(k) = A_c \sum_{n=0}^{N-1} w_n e^{\frac{j2\pi kn}{N}} \quad k=0, \dots, N-1$$

che nel caso di N pari ad una potenza di 2 risulta la IFFT (inversa della FFT, **trasformata di Fourier** rapida) del segnale modulato.

Quindi nell'OFDM la banda totale del canale viene suddivisa in sottobande (le portanti $\phi_n(t)$) opportunamente distanziate tra loro (ortogonalità). Il segnale da trasmettere viene associato ad alcune delle sottobande mentre altri segnali vengono associate ad altre sottobande libere ottenendo così un accesso multiplo al canale trasmissivo (OFDMA). La flessibilità della tecnologia permette di attivare o disattivare le sottobande a seconda dei casi e degli scenari trasmissivi.

La banda del segnale modulato OFDM è data da :

$$B = \frac{(N+1)}{T} = \frac{(N+1)}{(NT_s)} = \frac{(N+1)}{N} \frac{1}{T_s} \sim \frac{1}{T_s} \quad \text{se } N \gg 1$$

cioè per N molto maggiore di 1 la banda del segnale modulato è circa uguale all'inverso dell'intervallo di simbolo. L'**efficienza spettrale**

è pari a $\eta = R/B$ dove R rappresenta il *bit rate* e B la *banda*. **Essendo $R = (1/T_s) \log_2 M$, dove $M = 2^K$ rappresenta il numero di costellazioni della modulazione QAM ottenibile con K bit per simbolo, allora si avrà che $\eta = K$.**

Cioè l'efficienza spettrale aumenta se si incrementa il numero di costellazioni

della modulazione QAM ma al costo di una minore immunità agli errori. Le modulazioni impiegate sono la 4-QAM, la 16-QAM e la 64-QAM (Quadrature Amplitude Modulation) con efficienza spettrale crescente pari a 2, 4, 6 [bit/s/Hz] e vulnerabilità crescente in quanto i segnali modulati sono più vicini e quindi i disturbi tendono più facilmente a farli equivocare. OFDM risulta robusta rispetto alle interferenze intersimbolo (*ISI*) e rispetto a quelle da canale adiacente.

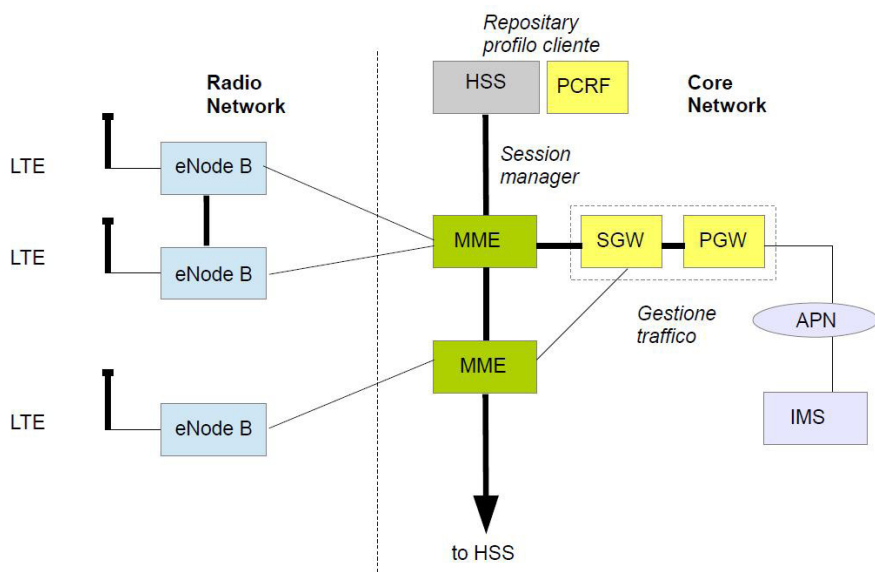
In uplink SC-FDMA è meno efficiente rispetto a OFDMA ma si adatta meglio all'amplificazione da parte dei terminali mobili.

La rete core di LTE è completamente a pacchetto IP e comprende una nuova core network denominata *EPC (Evolved Packet Core)*. L'accesso radio è costituito da un unico componente, l'*eNB (evolved NodeB)* con ruolo assimilabile a NodeB e RNC congiunti. L'eliminazione del nodo di controllo degli eNB (RNC), spostando le funzionalità sia sull'eNB, sia sui nodi di Core

Network consente una architettura più snella per conseguire l'obiettivo di minore latenza e maggior throughput.

La figura a seguire mostra l'architettura di rete LTE. La Core Network EPC presenta una completa separazione tra le funzioni di controllo e

quelle di trasporto.



servizio.

- *HSS (Home Subscriber Server)* contiene le informazioni del profilo utente e di posizione.

La terza ed ultima sostanziale differenza tra LTE e UMTS sta nell'utilizzo delle antenne. LTE utilizza la tecnologia MIMO (*Multiple-Input Multiple-Output*), cioè più antenne sia dal lato utente che dal lato della stazione base. Nelle altre tecniche, i cammini multipli dovuti a

Si descrivono ora i blocchi funzionali:

- *MME (Mobility Management Entity)* è un nodo di controllo, responsabile delle procedure di rete per l'autenticazione del terminale verso il data base di rete HSS (Home Subscriber Server), per l'instaurazione della connessione, per il mantenimento e il rilascio delle connessioni stesse e per gli aspetti di mobilità.
- *SGW (Serving Gateway)* è un nodo di trasporto, interlavora con il PGW (PDN Gateway) per il trasporto dei dati dell'utente gestendo la mobilità dell'utente tra accessi 3GPP e può cambiare in funzione dell'eNodeB di destinazione nella procedura di handover.
- *PGW (PDN Gateway)* è un nodo di trasporto assegnato dalla rete in dipendenza dal servizio richiesto e, anche in presenza di mobilità, rimane lo stesso per l'interconnessione alle reti dati PDN (Packet Data Network), sia interne al dominio dell'operatore sia esterne.
- *PCRF (Policy and Charging Rules Function)* è il responsabile per le policy (di QoS, charging, gating) su base utente e

riflessioni su palazzi o alberi (*multi-path fading*) erano un problema per la qualità del servizio. LTE invece sfrutta questa molteplicità di segnali andandoli a combinare tra loro per mezzo di più antenne (*tecniche di diversità*). Aggiungere antenne alla stazione radio base non è un grosso problema, mentre lo diventa sul terminale mobile per ragioni di spazio. Per questo dal lato utente si sfruttano antenne radio elettriche grazie ad array patch integrati. La tecnica MIMO consente di ottenere guadagni in termini di *capacità di canale* (bit rate massimo) pari al doppio rispetto ai precedenti sistemi.

Infine la tabella in basso riporta alcune cifre rispetto a quanto detto in precedenza sulla comparazione tra UMTS e LTE.

	UMTS	LTE
Accesso	W-CDMA	OFDMA
Efficienza spettrale	0.077	4
Commutazione	Pacchetto	Pacchetto IP
Download massimo(Mb/s)	0,384	300
Upload massimo(Mb/s)	0,384	75
Latenza(ms)	150	10

CONCLUSIONI

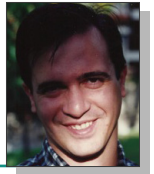
LTE benchè si ponga a metà strada tra terza e quarta generazione ha rappresentato comunque un passaggio storico dal punto di vista tecnologico per ciò che riguarda l'uso del protocollo IP per la totalità delle comunicazioni siano esse voce o dati. L'aumento della velocità di comunicazione (*bit rate*) e la limitata latenza lo rendono all'avanguardia per le applicazioni moderne su cellulare. Tuttavia questa corsa alla velocità vedrà un nuovo balzo in avanti grazie alla nuova versione di LTE denominata *Advanced* (LTE+) che promette 1Gb/s in termini di valori nominali di throughput.

RIFERIMENTI:

Easy LTE a cura di Paolo Semenzato.

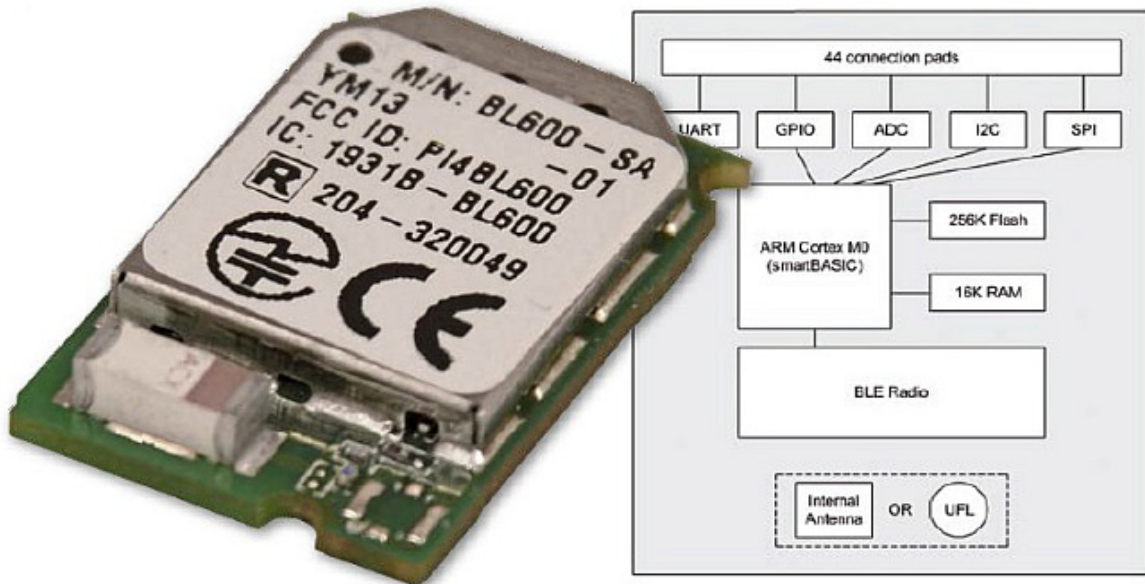
L'autore è a disposizione nei commenti per eventuali approfondimenti sul tema dell'Articolo.
Di seguito il link per accedere direttamente all'articolo sul Blog e partecipare alla discussione:
<http://it.emcelettronica.com/lte-scopriamo-la-telefonia-mobile-4g>

Termometro Bluetooth a basso consumo



di slovati

22 Settembre 2015



Grazie a questo progetto, chiunque sarà in grado di leggere in remoto il valore della temperatura ambientale, visualizzandolo comodamente su un qualunque smartphone. Il cuore del progetto è rappresentato dal modulo BL600-SA prodotto da Laird, un dispositivo Bluetooth 4.0 a bassa energia (BLE) racchiuso all'interno di un contenitore impermeabile e in grado di comunicare con uno smartphone Apple recente (dall'iPhone 4S in avanti), o con uno smartphone con sistema operativo Android 4.3 o superiore. Non sono richiesti altri componenti attivi, tutto il circuito è contenuto nel modulo, che è programmabile in... BASIC!

I Bluetooth mi ha sempre appassionato, e questo termometro rappresenta una buona scusa per utilizzare l'ultima versione dello standard, molto interessante in quanto la corrente richiesta è di gran lunga inferiore rispetto alle versioni precedenti (in particolare la 2.0 e la 1.0), con le quali il modulo BLE non è compatibile.

BLUETOOTH IN BASIC

Il linguaggio di programmazione utilizzato per il

modulo **BL600-SA** (la lettera 'A' indica che questa versione include un'antenna integrata) è lo smartBASIC, una versione del noto linguaggio di programmazione orientata agli eventi, in grado di gestire in modo semplice sia i sensori collegati direttamente al modulo che la trasmissione dei valori acquisiti verso un qualunque ricevitore **Bluetooth 4.0** (smartphone, tablet, computer, ecc.). E' quasi un gioco da ragazzi comunicare in modo wireless con piccoli dispositivi portatili,

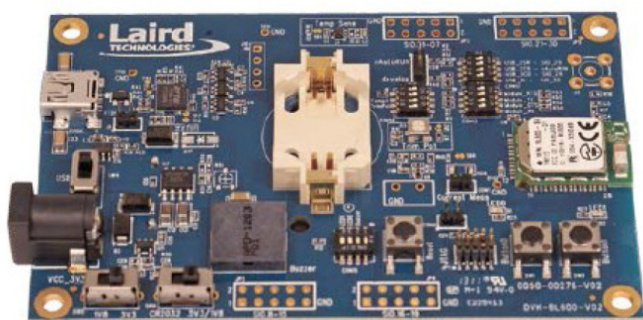
alimentati da pile di tipo AAA o a bottone.

Caratterizzato da un consumo di corrente estremamente contenuto, il modulo BL600-SA è basato sul chipset nRFS1822 prodotto da Nordic Semiconductor e include tutto l'hardware e il software richiesti per la comunicazione radio a 1 Mbps nella banda compresa tra 2,402 e 2,480 GHz.

Le caratteristiche principali del modulo sono le seguenti:

- interfacce UART, I2C, SPI;
- 28 segnali di ingresso/uscita di tipo general purpose (GPIO);
- 6 ingressi analogici (collegati a un ADC a 10-bit);
- assorbimento: 0,4 μ A nella modalità deep sleep, 5 μ A in modo stand-by, 10 mA durante la trasmissione;
- programmabile in linguaggio smartBASIC;
- dimensioni compatte: 19 x 12,5 x 3 mm;
- massima portata: fino a 20 m.

Impressionante, non vi pare? A tutto questo, aggiungiamo che il prezzo del modulo è molto competitivo. L'immagine seguente mostra il modulo alloggiato sopra l'evaluation board (SDK), anch'essa disponibile da Laird Technologies. Il produttore mira a utilizzare il modulo principalmente in applicazioni di telemetria in campo medico (misurazione della pressione sanguigna, del battito cardiaco, della temperatura corporea, ecc.), ma ovviamente non esistono limiti alle possibili applicazioni.



La prima volta che ho utilizzato il modulo, è bastata un'ora di lavoro per riuscire a trasmettere degli 1 e degli 0 inviati da un iPhone, in modo tale da far lampeggiare un LED. Grazie a questo incoraggiante avvio, ho finito per non contare più il tempo richiesto per completare il progetto del termometro ambientale. La documentazione che accompagna il componente è completa e facilmente comprensibile, così come il sito del produttore [4].

Occorre poi fare qualche considerazione sulla miniaturizzazione del componente: le sue dimensioni sono così ridotte che è molto difficile (anche se non impossibile) saldare il modulo a mano. Fortunatamente, Laird Technologies ha trovato un semplice trucco per tenere accuratamente in posizione il modulo (torneremo su questo punto più avanti, nell'articolo).

LAMPEGGIO LED

Grazie a questo primo, semplice programma per fare lampeggiare un LED, sono stato in grado di verificare che quando il LED è spento il modulo assorbe soltanto 5 μ A. Lo stesso principio verrà utilizzato per interrompere la funzionalità Bluetooth ad intervalli di tempo regolari, in modo tale da prolungare la durata delle batterie.

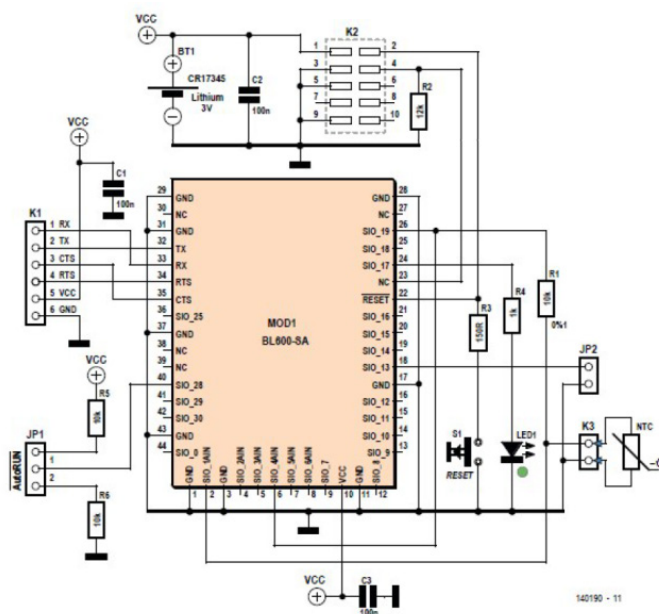
```
//*****
// Jennifer AUBINAIS 2014
// Test sleep with led
//*****
```

```
//*****
' LONG TIME : 2 seconds => led off
//*****
FUNCTION Func0()
Dim rc
' led off
rc = GPIOSetFunc(17,2,0)
TIMERSTART(1,2000,0)
ENDFUNC 1
```

resistenza, coefficiente termico α) fornita dal produttore dell'NTC [10], verrà eseguito sullo smartphone, tramite una app iOS o Android.

L'alimentazione viene fornita tramite il pin SIO_19 del modulo, in modo tale da ridurre l'assorbimento a $5\ \mu A$ in modo stand-by. Nella modalità deep-sleep, il consumo si riduce a $0,4\ \mu A$, ma si può risvegliare il modulo soltanto resettandolo oppure cambiando lo stato di un ingresso. Nella modalità stand-by, gli interrupt software consentono di ripristinare la comunicazione Bluetooth in breve tempo. E' stato scelto un periodo di sleep pari a 500 ms (configurabile nella funzione TIMERSTART). Una volta instaurata la comunicazione, l'assorbimento durante la trasmissione è pari a 10 mA. Il connettore K2 è stato previsto per possibili aggiornamenti futuri, tramite interfaccia JTAG, del firmware del modulo BL600. Il connettore K1 permette invece di programmare da un PC il BL600 tramite un'interfaccia seriale a 3,3 V. L'FT232R BOB offerto dall'e-shop Elektor [6] va benissimo per questo scopo.

L'analisi del circuito, visibile nell'immagine seguente, è abbastanza semplice. Per misurare la temperatura, viene utilizzato l'ADC del modulo BL600-SA (indicato con MOD1 nel diagramma), tramite una resistenza NTC collegata a K3. Il collegamento può essere anche remotizzato, la lunghezza del cavo a cui è collegata l'NTC non è critica. La misura della resistenza viene eseguita tramite un partitore di tensione, noto il valore di $R1 = 10\text{ k}\Omega\ 0,1\%$. L'equazione che permette di ricavare la temperatura in funzione del valore misurato dall'NTC è una funzione di tipo logaritmico che va al di là delle capacità di calcolo offerte dal modulo BL600. Questo calcolo, che richiede una tabella dati (temperatura,



Il LED svolge due funzioni distinte:

- debugging (con il jumper): il LED lampeg-

gia in modo continuativo, indicando le modalità stand-by e connessione del Bluetooth;

- normale (senza il jumper): il LED lampeggia brevemente durante l'inizializzazione per indicare che il termometro è stato avviato correttamente.

Durante i miei test iniziali, la trasmissione dati veniva troncata: dovevo fare una scelta tra risvegliare il modulo per un periodo maggiore oppure trasmettere una quantità inferiore di dati. La scelta è ricaduta sulla seconda opzione. Un esempio di trasmissione dati dal modulo è il seguente:

PW3012V853C433

PW: tensione di alimentazione (dalla batteria) da visualizzare nell'applicazione

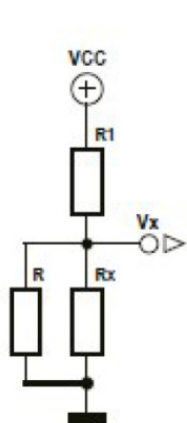
V: tensione sul partitore di tensione

C: tensione alla giunzione R1/NTC

A questi valori, come vedremo tra poco nella sezione relativa ai calcoli, corrisponde una temperatura di 24,3 °C.

I CALCOLI

Senza scendere troppo nei dettagli, possiamo subito notare come il valore del termistore possa essere calcolato misurando la tensione presente su un partitore.



$$V_x = V_{cc} \times \left(\frac{\frac{R R_x}{R + R_x}}{R1 + \frac{R R_x}{R + R_x}} \right)$$

$$\vdots$$

$$R_x = \frac{V_x R1 R}{V_{cc} R - V_x (R1 + R)}$$

Per quanto riguarda la temperatura, non utilizzerò la consueta equazione che lega il termistore alla temperatura tramite un coefficiente β :

$$R_{ctm} = R_{25} \times e^{\beta \times \left(\frac{1}{T} - \frac{1}{298} \right)}$$

Dove:

R_{25} : valore di resistenza alla temperatura di 25°C

β : coefficiente beta dell'NTC

T: temperatura in gradi kelvin

R_{ctn} : valore di NTC alla temperatura calcolata

Utilizzeremo invece un coefficiente α (%/T), il cui valore varia con la temperatura in base a una tabella fornita dal produttore del modulo [1], secondo la seguente formula:

$$R_T = R_{Tx} \cdot e^{\left[\frac{\alpha_x}{100} (Tx)^2 \cdot \left(\frac{1}{T} - \frac{1}{Tx} \right) \right]}$$

Dove:

R_T : valore di resistenza alla temperatura T

R_{Tx} : valore di resistenza alla temperatura x, come da tabella

T_x : temperatura in gradi kelvin inferiore alla temperatura misurata trovata in tabella

T: temperatura misurata, in gradi kelvin ($T_x < T < T_{x+1}$)

α_x : coefficiente termico alla temperatura T_x

Dalla stessa equazione possiamo ricavare la temperatura misurata T:

$$T = \frac{1}{\frac{100}{\alpha_x \cdot (Tx)^2} \cdot \ln \left(\frac{R_T}{R_{Tx}} \right) + \frac{1}{Tx}}$$

Table 1.

R/T No	4901	
T (°C)	$B_{25/100} = 3950 \text{ K}$	
	R_T/R_{25}	$\alpha \text{ (‰/K)}$
-30.0	16.915	6.1
-25.0	12.555	5.9
-20.0	9.4143	5.7
-15.0	7.1172	5.5
-10.0	5.4308	5.4
-5.0	4.1505	5.2
0.0	3.2014	5.0
5.0	2.5011	4.9
10.0	1.9691	4.7
15.0	1.5618	4.6
20.0	1.2474	4.5
25.0	1.0000	4.3
30.0	0.808	4.2
35.0	0.6569	4.1
40.0	0.5372	4.0

Come esempio di calcolo, consideriamo il seguente caso:

$$R_T = 10506,46 \, \Omega \text{ (} R_{\text{ctn}} \text{)}$$

$$\Rightarrow 12474 \, \Omega > R_T/T_{25} > 10000 \, \Omega$$

$$\Rightarrow \alpha x = 4,5$$

$$\Rightarrow T_x = 20^\circ\text{C} \text{ } 293,15 \text{ K}$$

$$\Rightarrow R_{Tx} = 12475 \, \Omega$$

$$T = \frac{1}{\frac{100}{4.5 \cdot (293.15)^2} \cdot \ln\left(\frac{10506.46}{12474}\right) + \frac{1}{293.15}}$$

$$T = 297.01 \text{ K}$$

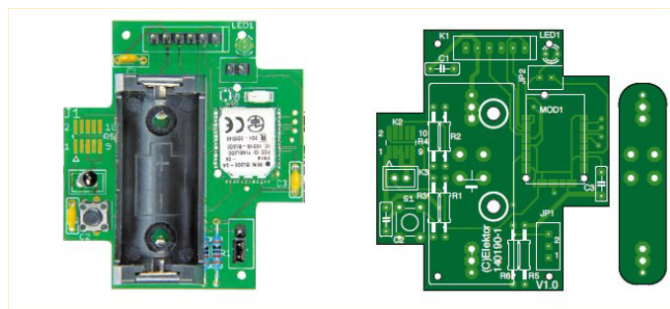
$$T = 297.01 - 273.15 = 23.86^\circ\text{C}$$

Se ci accontentiamo di un valore di accuratezza pari al 3%, per R_1 può essere utilizzata una comune resistenza con tolleranza dell'1% (anzi-

chè una con tolleranza pari a 0,1%).

ASSEMBLAGGIO

Una volta individuato un opportuno contenitore impermeabile, ho cominciato a disegnare il circuito stampato (visibile nell'immagine seguente) per la sonda di temperatura, nella speranza di non doverlo più modificare per almeno.. 10 anni! Ho inoltre creato una libreria Eagle per il modulo BL600 (disponibile nella sezione download dell'articolo [2]). Il contenitore di plastica è molto robusto e impermeabile, ma la sua struttura interna richiede al PCB di assumere una forma irregolare. Inizialmente ero un pò scettico riguardo l'effettivo assorbimento del modulo, per cui ho deciso di alimentarlo con una pila CR123 al posto di una normale CR2032.



La batteria viene inserita in un apposito adattatore, visibile nell'immagine qui sotto, che può essere inserito sul PCB tramite un connettore composto da una fila di pin. Per realizzare il connettore, occorre recuperare otto pin estraendoli dall'involucro plastico di un connettore SIL, saldarli sul PCB del modulo adattatore, e saldare infine il porta batteria CR123 sulla stessa board (prestando la dovuta attenzione alla polarità). Elektor, in ogni caso, fornisce il termometro Bluetooth sotto forma di modulo pre-assemblato, con il BL600 **già saldato** [3]. Chi volesse saldare il modulo BL600 da sè, dovrà anzitutto preparare uno stencil per applicare la pasta saldante. Anzitutto occorrerà posizionare tre

viti da 1,6 mm nei fori del PCB appositamente predisposti, e posizionare successivamente lo stencil prima di applicare la pasta saldante sul PCB. Per posizionare con accuratezza il modulo, inserire le tre viti da 1,6 mm **da sotto**, e successivamente inserire il modulo utilizzando le viti stesse come guide per un posizionamento corretto.

A questo punto non rimane che inserire la board nel forno di cottura.

Separare quindi otto pin header dal loro contenitore plastico, e saldarli direttamente sulla board. Tagliare poi la parte terminale di ogni pin, in modo tale da lasciare solo la parte con il socket: questi fori serviranno per posizionare sulla board l'adattatore sul quale è saldata la batteria. Non c'è molto altro dire riguardo ai restanti componenti, tutti di tipo passante. Una menzione merita R2, la cui funzione è soltanto quella di permettere gli aggiornamenti del firmware del modulo, e pertanto può essere omessa (così come il connettore K2).

Il jumper JP1 consente di selezionare due modalità distinte:

- *autorun* (posizione 1) - permette di eseguire il programma automaticamente, quando viene applicata l'alimentazione o a seguito di un reset;
- *command* (posizione 2) - permette di accedere al modulo tramite collegamento seriale e di inviare ad esso dei comandi AT. Ad esempio, il comando "AT+f 1" esegue una cancellazione della memoria di programma, esegue un restart del modulo, e permette di caricare il programma tramite l'interfaccia seriale.

Il jumper JP2 consente di entrare nella modalità debugging, e di redirigere i log del programma sulla porta seriale. Durante il funzionamento

normale, questa funzionalità è disabilitata.



ELENCO COMPONENTI

Resistenze

R1 = 10 kΩ 0,1%*

R2 = 12 kΩ*

R3 = 150 Ω

R4 = 1 kΩ

R5,R6 = 10 kΩ

termistore da 10 kΩ CTN B57891S103F8 (2112816)

Condensatori

C1,C2,C3 = 100 nF altezza 0,2"

Semiconduttori

LED1 = LED da 3mm

MOD1 = modulo BL600 (Laird Technologies)

Varie

JP1 = pinheader a 3-pin

JP2 = connettore a 2 vie es. Molex (9731148)

K1 = pinheader a 6-pin

S1 = pulsante miniatura (1555985)

porta batteria CR123A* (1650670)

16 pin socket per PCB

Contenitore di tipo Multicomp G302 (1094697)

PCB # 140190-1

Termometro Bluetooth già assemblato, disponibile su Elektor Store # 140190-91

*: si veda il testo. I numeri racchiusi tra parentesi tonde rappresentano i codici prodotto di Farnell/Newark

IL SOFTWARE

La programmazione del modulo BL600 è molto semplice, grazie al linguaggio smartBASIC di Laird Technologies. Si tratta di un BASIC molto simile a quello imparato a scuola, con semplici funzioni aritmetiche per facilitare l'elaborazione dei dati acquisiti. E' possibile creare sottoprogrammi e funzioni, gestire le funzionalità di I/O del modulo BL600, come anche le sue periferiche più complesse: I²C, SPI, ADC, e UART. Il programma può essere scritto utilizzando un qualunque editor di testo. Raccomandiamo come esempio Notepad++, anch'esso scaricabile dal sito di Laird Technologies. Il programma viene compilato e trasferito sul modulo utilizzando l'applicazione free UWTerminal.

Nella sezione download [2] è disponibile il codice sorgente del programma JATEMP, il quale indica come instaurare la comunicazione tra il modulo e uno smartphone, come inizializzare gli ADC, e come inviare al telefono i seguenti dati acquisiti: tensione della batteria, tensione sul partitore, tensione sul punto intermedio del partitore. I dati sono tutti trasmessi in formato stringa di caratteri: PW3012V853C433. L'app installata su iOS o Android esegue poi i necessari calcoli.

APP PER IOS E ANDROID

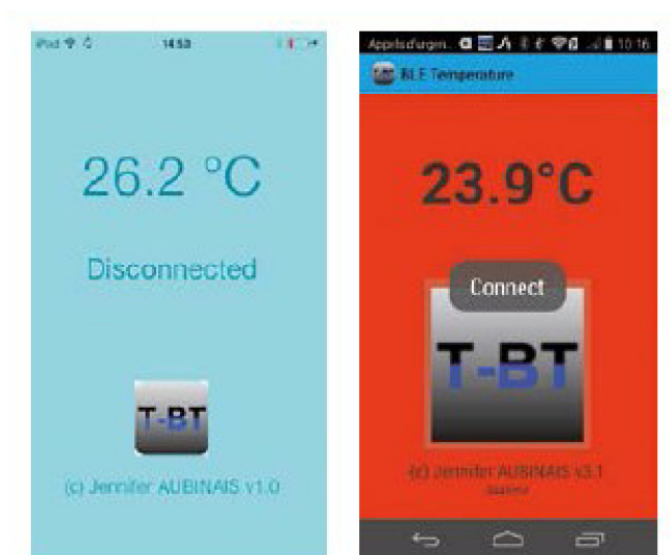
Laird Technologies permette di scaricare liberamente dal proprio sito il codice sorgente di un'applicazione seriale per il modulo BL600 su iOS. Poichè ho già una certa esperienza nello sviluppo su questo sistema operativo, ho preferito scrivere una mia applicazione [7] che visualiz-

za due tipi di informazione: temperatura e stato della connessione. Ho poi incorporato in essa i sorgenti del programma Serial BL600 disponibili sul sito Laird Technologies (UARTPeripheral.c e DataClass.c). Anzichè utilizzare un tasto per connettersi/disconnettersi con il modulo BL600, ho preferito creare un thread che controlla quando il termometro JATEMP si connette o si disconnette (l'operazione viene eseguita automaticamente dal termometro stesso).

Poichè non avevo alcuna esperienza di programmazione in Java, lo sviluppo dell'applicazione in ambiente Android mi ha creato non poche difficoltà. A questo proposito devo ringraziare il forum francese developpez.com, dove ho potuto trovare un aiuto soprattutto per creare un'interfaccia omogenea che si adattasse a tutti i dispositivi con sistema operativo Android. Ho quindi creato una mia applicazione Serial per il modulo BL600 (composta da una parte codice e da una parte interfaccia), molto semplice, in grado di collegarsi al termometro JATEMP. Il tasto "Scan and Connect" permette di ricevere dal termometro la stringa di caratteri in formato raw tramite un thread, esattamente come avviene sotto iOS. Il calcolo della temperatura, sia in iOS che in Android, avviene applicando le formule viste precedentemente. Lo sfondo dell'applicazione cambia in base alla temperatura misurata e, per evitare di rimanere indefinitamente collegati con il termometro, l'applicazione si interrompe automaticamente dopo 5 minuti, oppure quando viene premuto il tasto Home.

Per instaurare la comunicazione con lo smartphone e per collegarsi con il termometro la prima volta, **non è necessario eseguire alcuna operazione**, tutto avviene automaticamente. Nelle due immagini seguenti possiamo osservare come appaiono le applicazioni per iOS e

Android, ed il prototipo realizzato dall'autore.



IL FRIGORIFERO NON È UNA GABBIA DI FARADAY

Quando eseguii i miei primi test con il termometro era l'inizio di settembre 2014, e a Parigi si potevano apprezzare i benefici della cosiddetta "Estate Indiana": era perciò molto difficile misurare delle temperature inferiori a 20 °C per un periodo sufficientemente lungo. Decisi allora di collocare il termometro in una brocca di ghiaccio per circa un quarto d'ora, in modo da poter misurare anche le basse temperature. Mentre aspettavo che trascorresse questo lasso di tempo, andai a fare gli ultimi ritocchi alla mia applicazione Android. Con mia grande sorpresa, fui in grado di monitorare in tempo reale l'abbassamento della temperatura misurata dal termome-

tro. Come ulteriore verifica, decisi di collocare il termometro nella vaschetta del ghiaccio del mio vicino, verificando le variazioni di temperatura direttamente sul display dello smartphone. Potei subito constatare che la temperatura era scesa fino a circa 4 °C. Da questo semplice test potei dedurre che la porta del frigorifero permette il passaggio indisturbato del segnale Bluetooth (una caratteristica che i produttori di elettrodomestici dovrebbero menzionare nella scheda tecnica dei loro prodotti).

CONCLUSIONI

Spero con questo articolo di avervi convinto della validità e interesse del modulo BL600, i più arditi potranno anche cimentarsi nella saldatura a mano del modulo stesso.

La realizzazione di un'applicazione basata sul modulo BL600 presenta lo svantaggio offerto dalla miniaturizzazione dei contatti del componente. Tale svantaggio è però ampiamente compensato da un assemblaggio molto semplificato, e dalla semplicità di programmazione tramite il linguaggio smartBASIC. E questo è soltanto l'inizio..

Nel frattempo Laird Technologies sta proponendo sul mercato il modulo BL620, caratterizzato dalla possibilità di funzionare anche come master e di collegarsi ad altri moduli BL600. Questa caratteristica apre nuove prospettive che sia Elektor che l'autore vi inviteranno a scoprire nel prossimo futuro.

LINK

- [1] www.lairdtech.com/Products/Embedded-Solutions/Bluetooth-Radio-Modules/BL600-Series/
- [2] www.elektor-magazine.com/140190
- [3] www.elektor.com/bluetooth-thermometer
- [4] https://laird-ews-support.desk.com/?b_id=1945

- [5] www.youtube.com/watch?v=OYIKxtYwQiE
- [6] www.elektor.com/ft232r-usb-serial-bridge-bob-110553-91
- [7] <https://appsto.re/fr/XTwnV.i>
- [8] <https://play.google.com/store/apps/details?id=com.JA.bletemperature&hl=fr>
- [9] www.aubinais.net
- [10] -EPCOS NTC
www.epcos.com/inf/50/db/ntc_13/NTC_Leadled_disks_S891.pdf
www.epcos.com/blob/531152/download/2/pdf-standardizedrt.pdf

NTC Thermistors / General technical information:

www.physics.queensu.ca/~robbie/ENPH354/NTC-Therfnistors-Technote.pdf

Original article by Jennifer Aubinais courtesy of Elektor Electronics



L'autore è a disposizione nei commenti per eventuali approfondimenti sul tema dell'Articolo.
Di seguito il link per accedere direttamente all'articolo sul Blog e partecipare alla discussione:
<http://it.emcelettronica.com/termometro-bluetooth-a-basso-consumo>

SMA: gestiamo questi materiali “intelligenti” con Arduino



di Vincenzo
24 Settembre 2015



I fili a memoria di forma o Shape memory alloys (SMAs) sono materiali intelligenti che “ricordano” la loro forma originale. La loro caratteristica principale è quella di essere in grado di recuperare una forma preimpostata per effetto del semplice cambiamento di temperatura o dello stato di sollecitazione applicato. Gli SMAs sono molto usati come attuatori poiché sono materiali che cambiano forma, rigidità, posizione, frequenza naturale e altre caratteristiche meccaniche come risposta a gradienti di temperatura. La possibilità di utilizzo degli SMAs soprattutto come attuatori ha ampliato lo spettro di molti campi scientifici. Nell’ultimo decennio l’interesse scientifico nei confronti di questi materiali “intelligenti” è di molto aumentato a causa delle molteplici applicazioni e del loro basso costo. In questo articolo ne studieremo la struttura chimico-fisica e la risposta a cambiamenti di temperatura e a stress meccanici, in più, vedremo alcune particolari applicazioni ed esperimenti. Nella parte conclusiva, utilizzando Arduino, metteremo in piedi un esperimento DIY per meglio comprendere come si comportano gli SMAs.

LA STORIA

Gli SMAs sono delle leghe formate dall’unione di Nichel+Titanio, vengono solitamente chiamati “**Nitinol**”, “N” per Nickel, “t” per

Titanium e “nol” da Naval Ordnance Laboratory. La capacità di memoria di forma degli SMAs fu scoperta nel 1961 da William J. Buehler, un ricercatore del Naval Ordnance Laboratory

in *White Oak, Maryland*. La scoperta avvenne per caso. Durante una riunione di laboratorio, fu presentata una striscia di Nitinol che aveva subito diverse variazioni della sua forma, uno dei presenti, il *Dr. David S. Muzzey*, riscaldandola con la sua pipa si accorse che la striscia di Nitinol ritornò alla sua forma originale. Da quel momento fino ai giorni nostri, gli SMA hanno avuto un notevole successo e un largo spettro di applicazioni in ambito scientifico e non solo.

STRUTTURA CRISTALLINA

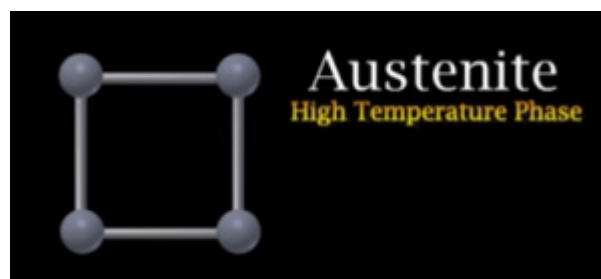
La maggior parte dei solidi ha un'unica struttura cristallina, ma il NiTi ne ha due! Il passaggio dall'una all'altra costituisce una vera e propria transizione di fase allo stato solido, analoga in molti aspetti alla transizione solido-liquido. Come in quest'ultimo caso, infatti, il passaggio tra due stati della materia è indotto dalla variazione di temperatura, dalla forza applicata, o da una combinazione dei due fattori.

La caratteristica distintiva di una transizione di fase è il brusco cambiamento di una o più proprietà fisiche. Le proprietà macroscopiche di un materiale infatti dipendono fortemente dalla struttura cristallina, pertanto una sua seppur lieve modificazione può portare a materiali anche molto diversi tra loro.

Un esempio familiare è quello del Carbonio: un cristallo di Carbonio con la struttura tetraedrica (*ogni atomo è al centro di un tetraedro ai cui vertici siedono altri atomi*) costituisce il *diamante*. Ma il Carbonio può ugualmente presentarsi nella forma cristallina della *grafite*, in cui gli atomi si dispongono su strati debolmente legati tra loro. Grafite e diamante hanno la stessa composizione chimica essendo fatti di soli atomi di carbonio. La loro struttura cristallina è però profondamente diversa.

Ogni struttura è stabile in un determinato intervallo di pressione e temperatura. Tuttavia, in genere una volta che il cristallo si è formato con una specifica forma cristallina (per esempio, il diamante si forma solo a pressioni molto elevate) è molto difficile fargli cambiare forma modificando semplicemente pressione e temperatura. Nel caso del Nitinol, invece, si può indurre il passaggio tra le due strutture variando la temperatura in un intervallo facilmente realizzabile in laboratorio. Le forme cristalline del Nitinol sono due:

- **Austenite**: è la fase stabile ad alta temperatura. È caratterizzata da un reticolo a simmetria cubica a corpo centrato (BCC), in cui gli atomi cioè occupano i vertici di due reticoli cubici compenetrati. Il materiale in questa fase è duro e difficilmente deformabile.

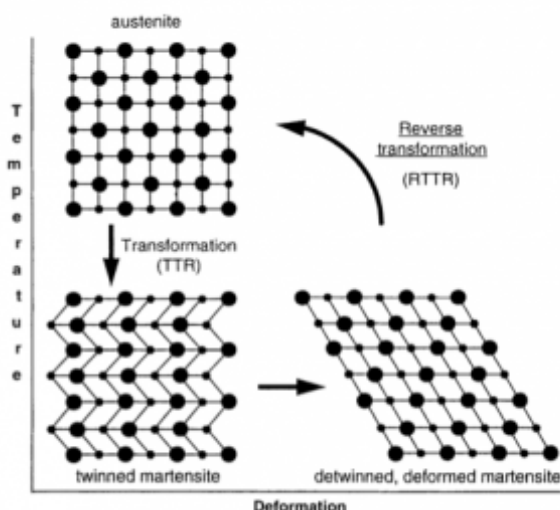


- **Martensite**: è la fase stabile a bassa temperatura. Ha una struttura cristallina monoclinica distorta, molto meno simmetrica del tipo BCC. È caratterizzata da grande flessibilità e dalla capacità di essere facilmente deformata senza che tuttavia tale deformazione sia permanente, infatti, questa fase si forma e si accomoda in forma *twinned*, ovvero speculare rispetto ad un piano ideale tra due celle, non creando difetti irreversibili nel reticolo cristallino. Questa struttura a twins è pertanto facilmente deformabile, aprendosi nella forma

detwinned come il mantice di una fisarmonica e, non avendo creato difetti di slittamento dei piani reticolari, è reversibile. Nella fase martensitica, perciò, il materiale sottoposto a uno sforzo meccanico è in grado di sopportare un alto grado di deformazione senza tuttavia rompere i legami chimici.



Poiché la trasformazione tra Austenite e Martensite è **reversibile**, alzando la temperatura il materiale esegue la trasformazione cristallina inversa e riprende la forma regolare e rigida della Austenite, indipendentemente dalla deformazione eventualmente subita nella fase Martensite. Si manifesta cioè l'**effetto di memoria di forma SME (Shape Memory Effect)**.



SUPERELASTICITÀ

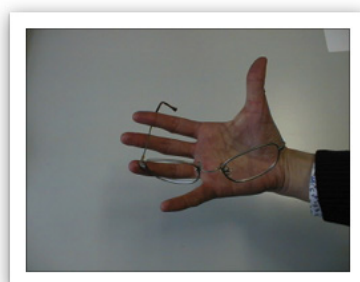
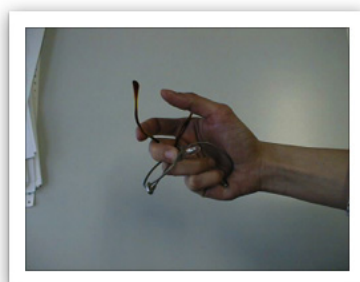
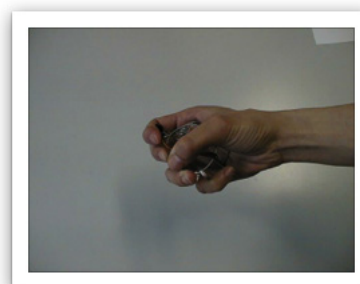
Oltre all'importante proprietà di "memoria" dei fili di Nitinol, un'altra importante caratteristica che li ha resi importanti dal punto di vista com-

merciale è sicuramente la **superelasticità**.

Tale proprietà si sviluppa quando la lega viene deformata, applicando un'appropriata forza, sopra la sua temperatura di trasformazione. Si genera in tal modo una Martensite indotta da sforzo (SIM) che si trova a temperatura maggiore del suo campo di esistenza: non appena lo sforzo viene rimosso essa si riconverte in Austenite indeformata. Questo fenomeno conferisce al materiale un'ottima elasticità.

Il fenomeno della superelasticità non è altro che un **effetto di memoria meccanica** del materiale: esso, sotto l'azione di uno stato di sollecitazione, assume una configurazione deformata, ben oltre il limite elastico, che può essere ripristinata togliendo lo stato di sollecitazione.

La superelasticità viene sfruttata per realizzazione di particolari montature per occhiali, come quelli mostrati qui di seguito:



in questo modo e' possibile realizzare montature per occhiali con eccezionali caratteristiche di resistenza alla deformazione e che, nel contempo, risultano particolarmente confortevoli da indossare.

COME SETTARE LA MEMORIA DI FORMA IN UN FILO DI NITINOL?

Non solo i "metalli intelligenti" possono ricordare la forma originaria, ma addirittura possono essere addestrati a "memorizzarne" una nuova. Per settare la forma che più preferiamo, dobbiamo tenere il materiale nella fase Austenite ad alte temperature. Possiamo creare uno "stampo" utilizzando delle viti per tenere fermo il filo su una superficie, ad una data forma. Il passo successivo è quello di riscaldare il tutto all'interno di un forno o con una fiamma continua. Ad alte temperature, il filo entra nella fase di Austenite, memorizzando la forma che abbiamo deciso di imprimergli. Quando lo raffreddiamo, immergendolo in acqua fredda, il filo entra nella fase Martensite ma esso lo fa senza cambiare la forma, poichè è ancora vincolato al nostro "stampo".

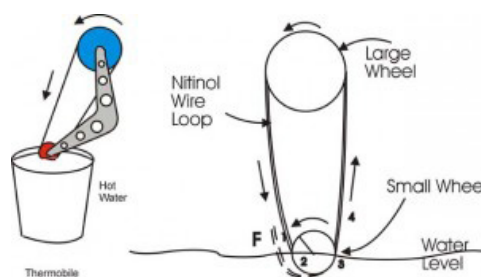
Una volta rimosso il filo di Nitinol dallo stampo, possiamo deformarlo come preferiamo. Questo provoca un cambiamento di forma permanente. Se successivamente riscaldiamo il filo, esso passa dalla fase Martensite alla Austenite, ritornando indietro, alla forma originale. Così facendo, osserviamo la sua capacità di memoria. Una volta fissata la forma, possiamo ripetere per quante volte vogliamo gli step di deformazione e successivo riscaldamento del filo.

MOTORI TERMICI CON I FILI DI NITINOL: IL THERMOBILE

Il Thermobile è un modello di motore termico

che usa un di filo di Nitinol per generare potenza. Il filo di Nitinol è posto su due ruote libere rotanti. Una di queste due ruote viene tenuta a diretto contatto con l'acqua calda (lato caldo) mentre l'altra è a diretto contatto con l'aria fredda (lato freddo). Il filo di nitinol, prima di essere montato sulla struttura, ha subito un trattamento di impressione della forma. Per ogni giro, ogni volta che il filo passa all'interno del recipiente con acqua calda superando la sua temperatura di transizione, si "ricorda" della sua forma originale e tenta di raddrizzarsi.

Se osserviamo la figura seguente, possiamo dire che nella posizione 1 il filo è "freddo", quando passa dalla posizione 1 alla 2, è piegato attorno alla rotella di ottone ed entra nell'acqua calda. Quando il filo passa dalla posizione 2 alla 3 viene superata la temperatura di transizione e questo tenta di raddrizzarsi. Dalla posizione 3 alla 4 il filo si è raddrizzato e ha permesso alla ruota di ruotare. Quando passa dalla posizione 4 alla 1, attraversando la ruota più grande, il filo, ha il tempo sufficiente per raffreddarsi scendendo al di sotto della sua temperatura di transizione, pronto per un altro giro.



Modello reale del Thermobile

Riassumendo, la differenza di temperatura provoca la contrazione del filo nel lato della ruota posta in acqua calda. Il filo si rilassa quando si trova a contatto con l'aria fredda. Ciò provoca una forza meccanica che garantisce la rotazione continua delle ruote. Lo stesso risultato può essere ottenuto utilizzando, al posto del recipiente con acqua calda, un raggio di sole focalizzato da una lente in un punto del filo di Nitinol. In alcuni casi è necessario fornire un avviamento alla ruota applicando noi stessi una forza esterna.

Nel 1982 Innovative Technologies International (ITI) costruì un motore che usava fili da 0.558 mm di diametro. Il motore fu testato usando acqua calda a 55°C e aria a 25°C. Il motore riuscì a raggiungere un'avelocità di 270 rpm (giri per minuto) e continuò a funzionare per un anno e mezzo.

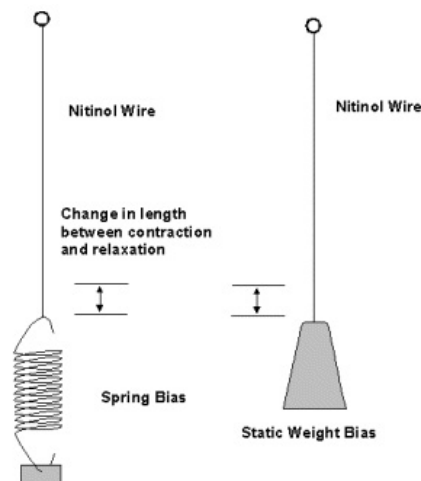
ESPERIMENTO 1

Il filo di Nitinol genera una forza per riacquisire la forma di 22.000 psi (**pound per square inch** o *libbre per pollice quadrato*). Considereremo fili da 0.006 pollici di diametro e da 0.015 pollici di diametro. Il primo ha una forza di contrazione maggiore, rispetto al secondo e può contrarsi fino all' 8%/10% della sua lunghezza.

Contrazione e rilassamento dipendono esclusivamente dalla temperatura del filo di Nitinol. Per riscaldare e raffreddare il filo può essere utilizzato qualsiasi metodo. Un modo solitamente usato per riscaldare il filo di Nitinol è quello di farlo attraversare da **corrente elettrica**. Il filo di Nitinol rappresenta una resistenza di circa **1.25 ohm** per pollice per il filo da 0.006 pollici di diametro. La resistenza del filo alla corrente elettrica genera calore sufficiente (**riscaldamento ohmico** detto anche "**effetto Joule**") da por-

tarlo alla sua temperatura di transizione in un tempo che è funzione della resistenza e della corrente.

Potremmo applicare al filo una controforza nella direzione opposta alla sua contrazione. La controforza distende il filo alla sua forma originale durante la fase di bassa temperatura. Questa controforza è chiamata **forza di bias**.



Se il filo di Nitinol è portato alla sua temperatura di transizione senza forza di bias esso si contrarrà, ma, quando si raffredda, esso non ritornerà alla sua lunghezza originale proprio perchè non ha un contro peso che gli consenta di distendersi. Di conseguenza, senza una forza di bias, quando il filo viene riscaldato nuovamente non avrà luogo nessuna ulteriore contrazione. La figura di sopra mostra due metodi per applicare una forza di bias: nel primo caso una molla, nel secondo caso un peso statico.

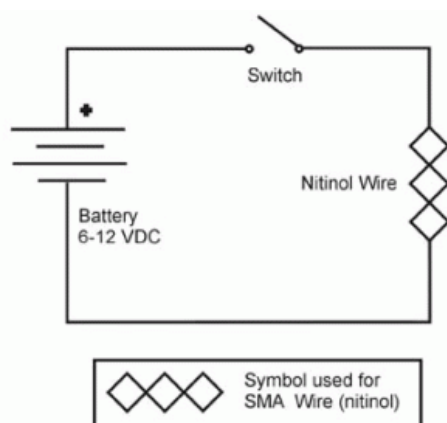
La velocità e la forza di contrazione del filo dipendono da quanto alta e da quanto velocemente viene incrementata la temperatura del filo. Ad esempio, 400mA di corrente elettrica attraverso il filo di nitinol di 0.006 pollici di diametro produrrà un sollevamento massimo di 11 grammi raggiungendo la completa contrazione in 1 secondo. Il tempo di reazione può essere più veloce, dell'ordine del millisecondo.

Per raggiungere queste alte correnti sono usati impulsi di breve durata. Nel fare questo bisogna considerare la massa e la velocità del materiale da spostare.

Quanto più velocemente vogliamo muovere una data massa, tanto maggiore sarà l'inerzia che dovrà essere superata. Se l'inerzia diventa maggiore di quelle sostenibili dal filo, esso potrà spezzarsi.

RISCALDAMENTO ATTRAVERSO CORRENTE ELETTRICA DIRETTA

Il filo di Nitinol può essere "attivato" (cioè portato alla sua temperatura di transizione) con una bassa tensione di alimentazione in DC (6-12Volt). Possiamo costruire un semplice circuito utilizzando una batteria, un interruttore e un pezzo di filo di Nitinol, come mostrato di seguito:

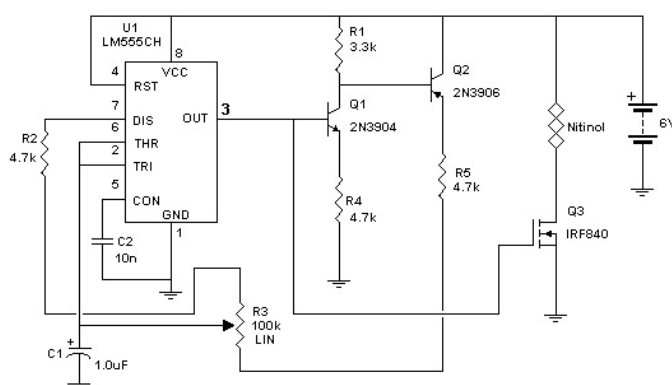


Quando si porta il filo nel suo "stato di attivazione" utilizzando corrente continua è importante non surriscaldarlo, altrimenti si correrebbe il rischio di apportare dei danni permanenti alle sue proprietà fondamentali. È importante osservare che **una corrente continua non riesce a surriscaldare il filo di Nitinol in maniera uniforme. Il metodo migliore è riscaldare il filo con un segnale PWM.** Usando un segnale PWM possiamo cambiare il duty cycle proporzionalmente alla quantità di contrazione che vogliamo ottenere. Questo tipo di controllo ci permette di "at-

tivare" il filo di Nitinol per lunghi periodi di tempo senza causare danni causati dal sovrariscaldamento della struttura cristallina del Nitinol.

Possiamo realizzare un circuito per generare un segnale PWM.

Il circuito sarà costituito da un timer 555. Utilizzando due transistor Q1 e Q2 e un potenziometro con il timer 555 possiamo creare un'uscita con una relativa frequenza costante con un duty cycle variabile. Ecco mostrato di seguito il circuito:



Per capire come funziona nel dettaglio un timer 555 rimando all'articolo scritto da Marven.

APPLICAZIONI FUTURE

Oggi, le proprietà di memoria di forma hanno reso il Nitinol un materiale ideale in ambiti molto diversi tra loro, dalle missioni spaziali alle decorazioni floreali (farfalle e libellule animate), dagli stent vascolari utilizzati per garantire il flusso sanguigno nelle arterie otturate, agli attuatori (tendini artificiali) per microrobot operanti tramite effetto Joule. L'estrema flessibilità del nitinol trova impiego anche nelle antenne dei cellulari, negli apparecchi ortodontici, nelle montature di occhiali, nelle placche per saldare fratture ossee, come attuatori nelle protesi biomedicali. Per quanto riguarda il futuro, la ricerca si sta impegnando sull'includere motori, nelle macchine, negli aerei e nei generatori elettrici, che

utilizzano l'energia meccanica derivante dalle trasformazioni di forma. È già in corso lo studio per l'applicazione dei fili di Nitinol nelle carrozzerie delle auto:

“ proviamo a immaginare che deformando la carrozzeria di un'auto dopo un incidente, sia possibile far ritornare la carrozzeria nella sua forma originale semplicemente riscaldando la superficie danneggiata.

Si pensa che, quanto sopra detto, sarà una tecnologia alla portata di tutti tra circa un decennio. Altre applicazioni in campo automotive sono previste per sistemi di raffreddamento del motore, controlli di lubrificazione e per ridurre il flusso d'aria attraverso il radiatore all'avvio quando il motore è freddo e quindi a ridurre il consumo di carburante e le emissioni.

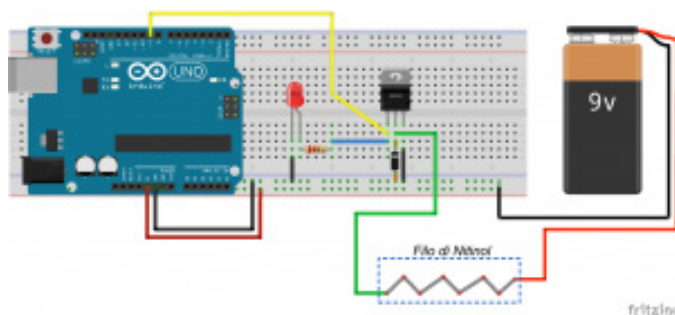
ESPERIMENTO CON ARDUINO

Materiale occorrente:

- Arduino Uno
- Potenzimetro 10kΩ
- Tip 120
- breadboard
- jumpers
- connettori a coccodrillo
- batteria 9V
- Led
- Diodo
- filo di Nitinol (acquistabile su [KRL](https://www.krl.it))

Di seguito vedremo due diversi esperimenti realizzati con Arduino e semplicemente riproducibili. Questi esperimenti ci permettono di capire meglio il comportamento degli SMAs

SCHEMA 1



Schema 1

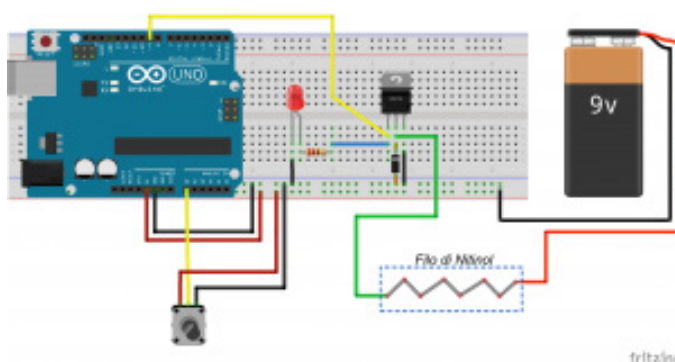
Il codice da caricare è il seguente:

```
int transistorPin = 9;

void setup() {
  pinMode(transistorPin, OUTPUT);
}

void loop() {
  digitalWrite(transistorPin, HIGH);
  delay(8000);
  digitalWrite(transistorPin, LOW);
  delay(8000);
}
```

SCHEMA 2



Schema 2

Il codice da caricare è il seguente:

```

int transistorPin = 9;

void setup() {
  pinMode(transistorPin, OUTPUT);
}

void loop() {
  int sensorValue = analogRead(A0);
  int outputValue = map(sensorValue, 0, 1023, 0,
255);
  analogWrite(transistorPin, outputValue);
}

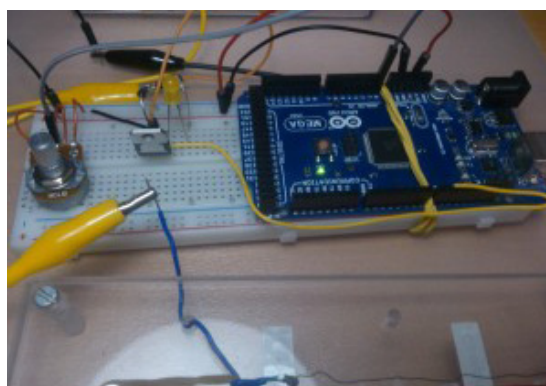
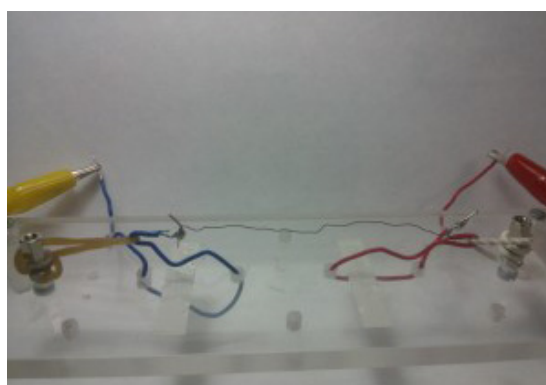
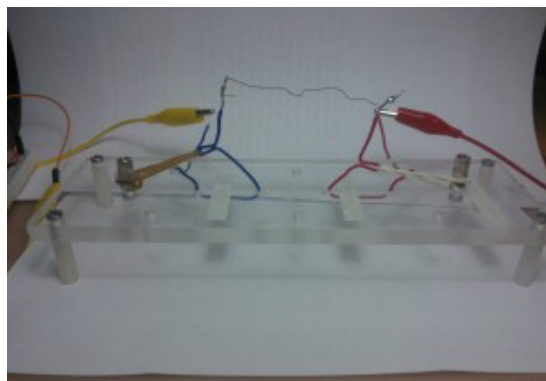
```

Per quanto riguarda il codice riferito allo schema 1, è stato utilizzato l'esempio del Blink Led di Arduino con un delay di 8 secondi, utile a permettere il raffreddamento del filo di Nitinol permettendone la successiva contrazione. Ogni 8 secondi dal pin 9 di Arduino vengono erogati 5V di tensione che portano in conduzione il TIP120 il quale amplifica la corrente destinata ad attraversare il filo di Nitinol. A causa della ridotta corrente di uscita dal pin 9 (40mA) è stato inserito un transistor NPN modello TIP120. Per portare a buon fine l'esperimento è necessaria un'alimentazione esterna, nel caso specifico è stata utilizzata una batteria da 9V ma il consiglio è quello di utilizzare un alimentatore da laboratorio. Ad ogni contrazione del filo, pilotata dal pin 9, vengono assorbiti circa 800mA. **La corrente assorbita è proporzionale alla resistività intrinseca del filo**, infatti, a seconda delle caratteristiche tecniche del filo di Nitinol, come ad esempio il diametro, la corrente assorbita che permette l'ingresso in fase Austenite sarà diversa.

Per quanto concerne lo Schema 2, è stato aggiunto un potenziometro in modo da regolare il duty cycle del segnale PWM in uscita dal pin 9. Nel codice relativo al secondo schema il pin 9 non si occupa più di erogare una tensione conti-

nua di uscita di 5V ma è stato programmato per erogare un segnale PWM.

Qui di seguito alcune foto della fase di test in laboratorio:



Circuito reale

Per eseguire i test, basandomi sui precedenti schemi, ho realizzato una particolare struttura che mi consentisse di mantenere teso il filo di Nitinol. La struttura è caratterizzata da una base in plexiglas forata in cui ho inserito due normallissimi fili con anima rigida a cui ho saldato il filo di Nitinol. Per garantire la tensione al filo di Nitinol ho utilizzato due elastici ancorati ai fili con anima rigida.

Per la realizzazione del struttura ho preso ispirazione da un esperimento realizzato da un certo Anton e pubblicato su youtube che è possibile vedere cliccando sul seguente link.

Durante la realizzazione dell'esperimento vi è stato un interessante scambio di informazioni via e-mail con Anton, dove ci si confrontava sui limiti e le possibili applicazioni dei fili di Nitinol.

Questa è un ulteriore prova di come l'open source avvicini appassionati con interessi comuni e ne nasca un grande confronto e una crescita reciproca.

AVVERTIMENTI E PRECAUZIONI

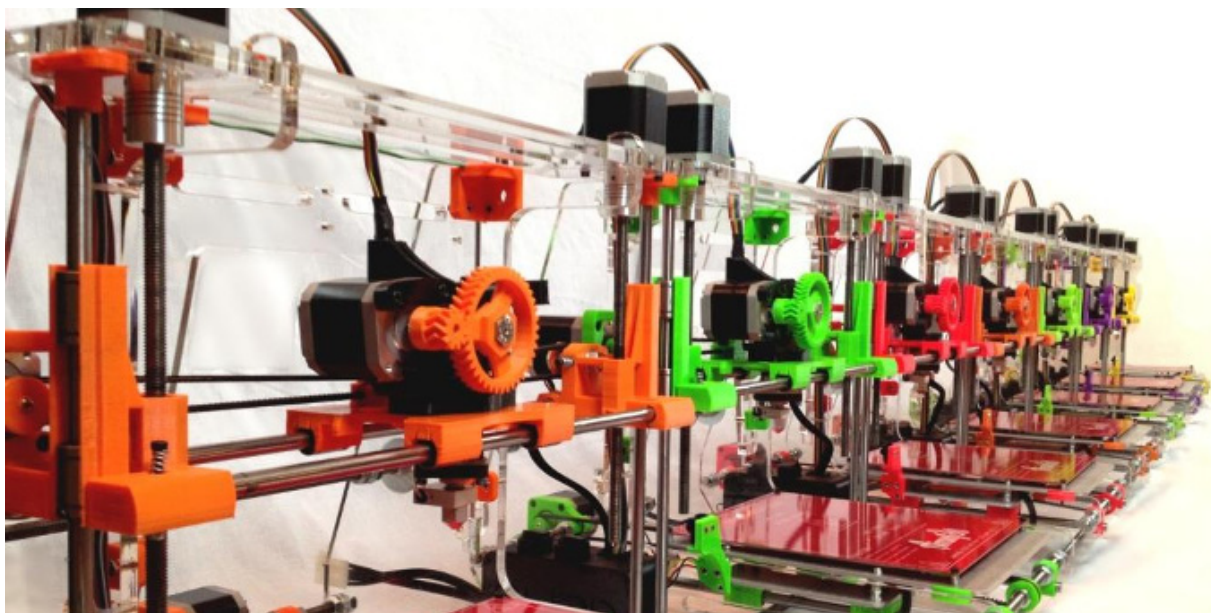
È doveroso consigliare di prestare la massima attenzione durante l' esperimento. I possibili pericoli possono derivare dalla diminuzione del delay (nel codice dello schema 1) o dalla diminuzione drastica del valore del potenziometro (nello schema 2). Entrambi i casi, infatti, portano il filo di Nitinol a un pesante stress dovuto alla tensione continua che lo attraversa per un tempo prolungato provocandone il surriscaldamento e la conseguente incandescenza. Prestiamo, inoltre, molta attenzione a non toccare il circuito durante il funzionamento a causa delle alte correnti in gioco. Il TIP120 raggiungerà temperature abbastanza elevate, il consiglio è quello di installare un dissipatore per evitare di danneggiare il dispositivo a lungo termine.

L'autore è a disposizione nei commenti per eventuali approfondimenti sul tema dell'Articolo.
Di seguito il link per accedere direttamente all'articolo sul Blog e partecipare alla discussione:
<http://it.emcelettronica.com/smas-gestiamo-questi-materiali-intelligenti-con-arduino>

Cosa serve per costruire una stampante 3D?



di [Piero Boccadoro](#)
29 Settembre 2015



La stampa 3D è un argomento che appassiona molti. Questo blog è seguito da una serie di maker e professionisti che sulla stampa 3D stanno lavorando tantissimo. Abbiamo gli strumenti, abbiamo le competenze e abbiamo la passione e questo è tutto ciò che serve per spiegare, a chi invece sta imparando, come fare per creare il proprio modello. In questo articolo analizzeremo la stampa 3D dal punto di vista della realizzazione dello strumento che vi accompagnerà nella prototipazione di pezzi e nella realizzazione dei vostri oggetti. Buona lettura.

Per creare una stampante 3D occorrono buone competenze di meccanica e di elettronica; servono anche tanti strumenti, dai cacciaviti alle pinze. Ma ciò che serve più di tutto è la passione, che vi guiderà anche quando non avrete una visione complessiva o particolarmente competente di ciò che state provando a realizzare.

Questo articolo nasce proprio con questo intento, darvi tutti gli strumenti, almeno teorici, per

poter mettere in pratica la vostra voglia di fare e creare il vostro modello di stampante.

In verità, nel recente passato, abbiamo parlato di diversi aspetti della [stampa 3D](#) e di alcune stampanti in [particolare](#), riguardo esperienze di [montaggio](#) ma abbiamo anche trattato le [prospettive](#) di questa tecnica e le sue frontiere, vicine e lontane. Oggi però proviamo a fare un passo ulteriore e a trattare questo argomento in maniera un po' più particolareggiata, cambian-

do il punto di vista.

SPECIFICHE TECNICHE

Come per tutti i progetti, la domanda fondamentale che dovete farvi è: **cosa voglio realizzare?**

La stampante 3D non è una macchina che può realizzare qualunque pezzo ma una soluzione ottimizzata in grado di gestire la creazione e la lavorazione di elementi e prototipi che prima di tutto hanno una loro dimensione e poi una precisione realizzativa.

In via preliminare, vi serve farvi alcune domande:

- **Quale dev'essere il suo ingombro fisico complessivo?**
- **Quale la dimensione massima dell'oggetto che posso realizzare?**
- **Con quale precisione voglio essere in grado di realizzare il mio oggetto?**
- **Quanto velocemente voglio che la macchina sia in grado di stampare?**
- **Voglio poterla connettere soltanto al computer oppure voglio renderla autonoma?**

E di conseguenza, poi, renderla autonoma vuol dire **in grado di connettersi ad altri dispositivi** in modalità wireless oppure voglio utilizzare un supporto di memoria esterno da cui estrarre informazioni su cosa stampare?

Ciascuna domanda tra quelle che abbiamo elencato comporta delle conseguenze e delle scelte progettuali. **Partiamo dalla prima: quanto spazio abbiamo a disposizione?** In che ambiente andremo ad inserire la nostra macchina? Abbiamo a disposizione un'officina? Un capannone? Una semplice stanza in un appartamento? Non soltanto le dimensioni fisiche saranno un parametro fortemente inficiato da queste risposte ma anche, per esempio, la ge-

stione di alcuni materiali che necessitano di una certa attenzione all'emissione di fumi, come per esempio l'ABS.

La seconda domanda riguarda, invece, **il dettaglio sull'oggetto finito**. Naturalmente l'area di stampa sarà interna alla dimensione fisica della stampante e pertanto saranno necessariamente più piccoli. Ma di quanto? Com'è fatta la struttura? Che tipo di oggetti dobbiamo realizzare? Vogliamo un macchinario professionale oppure da utilizzare per hobby?

La precisione realizzativa, ovvero il più piccolo dettaglio che siamo in grado di realizzare, è uno dei parametri fondamentali. Una cifra di merito estremamente valida è pari a **60 micron**. Questa è la dimensione che oggi rappresenta il miglior risultato ottenibile con stampanti di tipo amatoriale e anche con le professionali ma di fascia più bassa. Per intenderci, stiamo parlando di tutti i modelli di stampante 3D a deposizione di materiale fuso che non costano oltre i 4.000 €. È una risoluzione di tutto rispetto, intendiamoci, sfida anche l'effettiva necessità in ambito professionale, soprattutto in alcuni casi. Naturalmente si può fare anche di peggio, con valori che superano anche i 200 micron, ovvero 0,2 mm. Diciamo che riuscire a fare stampe con valori compresi **fra i 60 micron ed i 400 micron** rappresenta l'ottimo per un buon numero di applicazioni e una vasta gamma di casi. Questo è quello a cui bisognerebbe puntare per ottenere un buon modello di stampante.

La velocità di stampa è un parametro molto delicato. L'FDM è una tecnica molto veloce di per sé ma naturalmente una persona inesperta e che non opera in questo settore, non occupandosi di prototipazione rapida, non è necessariamente in grado di capire se il tempo di un'ora sia eccessivo oppure commisurato al tipo di pezzo

o al tipo di precisione. Diciamo che

“ quello che vogliamo è che la stampante operi velocemente ma naturalmente vogliamo anche che sia precisa e queste due specifiche, come ogni problema ingegneristico, sono sempre in controtendenza tra loro.

Data la delicatezza della macchina, la specificità delle sollecitazioni meccaniche, la velocità alla quale è il caso di far operare motorini, come vedremo più tardi, non è assolutamente mai il caso di superare una **velocità di 100 mm/s**.

Ovviamente esistono anche altre velocità di riferimento perché quella di cui abbiamo parlato fino a questo momento riguarda la velocità della stampa, durante la stampa. Ma **la macchina esegue dei movimenti anche quando non sta estrudendo** e allora viene naturale chiedersi: **in questi casi la velocità massima accettabile è la stessa?** O posso andare più veloce? O devo andare più lento?

In linea generale sappiate comunque che **più veloce andrà la vostra macchina più rischi avrete sia nella stabilità sia nelle vibrazioni**, entrambi parametri che inficeranno o potrebbero inficiare la qualità e la resa della stampa, sia nella taratura, che nella migliore delle ipotesi dovrete rifare più spesso, ovvero dopo meno ore di stampa, **fino ad arrivare addirittura al 50% di autonomia in meno** tra una stampa e la successiva.

Le esigenze di **connettività** stanno diventando sempre più evidenti dal momento che i modelli commerciali prevedono stampanti in grado di

effettuare **connessioni Wi-Fi, Bluetooth** o più genericamente wireless. Il fatto che una stampante possa essere connessa a dispositivi in rete è certamente una comodità ma rende il progetto più laborioso, complesso e da certi punti di vista anche più debole per via del fatto che dovrete occuparvi di problemi di interferenza di varia natura.

L'altra possibilità è quella di operare offline, con collegamento via cavo diretto al computer oppure con lettura di una scheda di memoria sulla quale il programma che effettua lo slicing memorizza il G-Code.

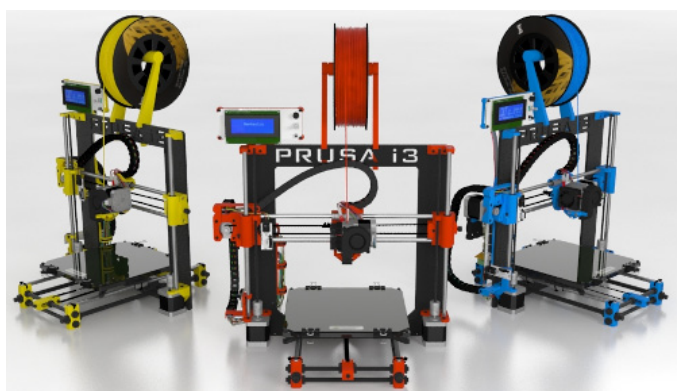
Si tratta di una scelta che implica la giunta o la sottrazione dal progetto hardware di un pezzo, ovvero quello relativo alla connettività che scegliete. E pertanto dovrete decidere prima quale meccanismo di connessione utilizzare e poi lavorare sull'elettronica e sulla programmazione per garantire il funzionamento.

STRUTTURA E MECCANICA

La stampante 3D non è altro che una macchina a controllo numerico (CNC, ovvero Computer Numerical Control) in grado di movimentare degli elementi altrimenti statici che servono a determinate operazioni. Si possono movimentare degli assi, in maniera tale che le coordinate cartesiane corrispondano a posizioni in uno spazio.

Ma di che tipo di spazio stiamo parlando? **Esistono diverse filosofie di gestione dello spazio dell'area di stampa e diverse geometrie: esistono le strutture a portale, la struttura delta, e altre ancora.** La scelta influenza il meccanismo di funzionamento, la disposizione degli elementi sugli assi e tanti altri parametri.

Prendiamo, per esempio, il modello Prusa, celeberrimo e di grandissimo successo.



La struttura di questa soluzione prevede la movimentazione del piano lungo un asse e dell'estrusore con due motorini. Questo non è soltanto un fatto logistico oppure di organizzazione ma ha un effetto diretto sulla scrittura del firmware e sulla gestione dei componenti e non solo.

Se consideriamo, invece, adesso, un'altra struttura, il modello delta, ci rendiamo conto che la filosofia di approccio al problema è radicalmente diversa.



In questa struttura non si ragiona più tanto in coordinate cartesiane ma piuttosto si cerca di gestire il problema del posizionamento compensando la posizione dell'estrusore tramite la movimentazione di un motorino in funzione della posizione degli altri. Tutti e tre salgono rispettivamente sui loro assi ed è dalla reciproca quota che complessivamente si ottiene il posizionamento dell'estrusore.

Altra cosa che cambia in questo approccio è che **l'elemento più grande che possiamo pensare di realizzare non è un cubo perchè la base non è un quadrato ma un triangolo**. Questo può sembrare banale ma bisogna pensarci perchè l'utilizzo ottimale del piano di stampa inficia la dimensione massima dell'oggetto che possiamo ottenere e questo, ancora una volta, è uno dei parametri che abbiamo tenuto in considerazione fin dall'inizio.

Questo secondo tipo di struttura presenta alcuni vantaggi tutt'altro che trascurabili, in particolare il fatto che la dimensione verticale è particolarmente privilegiata rispetto alle altre. Questo significa che possiamo scegliere questa struttura per questo approccio al problema quando vogliamo realizzare pezzi effettivamente molto "sproporzionati".

Tornando a quello che dicevamo prima, **si possono movimentare elementi, tra cui il filamento che deve essere fuso**. Ecco per quale motivo servirà sicuramente un motorino dedicato al trascinamento del filamento. Anche questo si propone come un problema di azionamento meccanico e di funzionamento e se ci pensate è ovvio: **a che velocità deve essere tirato il filamento?** Come impatta questo parametro con la resa della stampa?

La risposta alla prima domanda dipende dallo spessore del filamento, dal motorino e dalla ve-

locità di stampa per cui facciamo in modo da supporre che si possa rispondere in un secondo momento.

Per quanto riguarda la seconda domanda, invece, a questa possiamo rispondere subito: il flusso deve essere costante, la velocità con la quale il filamento attraversa la regione calda dipende dal primo parametro ma fondamentalmente il problema diventa quello di fornire al piatto di stampa una quantità di materiale che sia di dimensione e diametro costante in funzione di quello che effettivamente abbiamo inserito come parametro realizzativo della stampa stessa, ovvero lo spessore dello strato, del layer.

È tutto chiaro?

Se sì, possiamo passare oltre perché a questo punto dobbiamo parlare di un argomento abbastanza delicato che ha ancora una volta una grande influenza sulla struttura meccanica dell'intera macchina: **i profilati e la struttura in che materiale devono essere realizzati?**

Ci sono soluzioni che prevedono l'acciaio, altre che prevedono il legno. Quale soluzione adottare allora? Qual è la migliore?

Non si tratta di una risposta banale e per questo non merita di essere trattato in maniera sintetica. Questo punto, infatti, è fondamentale per la stabilità della macchina soprattutto mentre opera. Il legno è un materiale flessibile, che vibra e non assorbe le vibrazioni come molti potrebbero pensare. Sebbene ci siano, quindi, diversi maker che, per estetica o per scelta, optano per l'utilizzo del legno per tutta la parte dell'intelaiatura, buona norma vuole che si utilizzi materiale ferroso.

Per assorbire le vibrazioni della macchina durante le operazioni **si possono usare strutture**

di base dotate di molle, di fatto replicando la struttura delle case antisismiche che assorbono le vibrazioni piuttosto che trasferirle alle fondamenta o peggio ancora alle pareti.

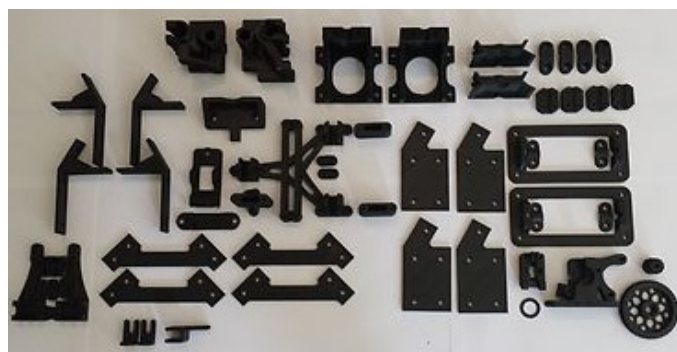
Per progettare una stampante 3D occorreranno una serie di pezzi meccanici di natura abbastanza variegata perché dovrete certamente realizzare la struttura, un'intelaiatura rigida che naturalmente richiederà dei pezzi statici che non devono essere movimentati ma che devono necessariamente sorreggere il tutto.

Tra questi elementi ci sono proprio i profilati di cui abbiamo parlato. E a questo punto sapete anche di che dimensioni sono, visto che avete risposto alle domande preliminari. Giusto?

Vi serviranno delle aste attorno a cui far spostare dei pezzi, così come dei cuscinetti e dei profilati attorno ai quali assemblare il tutto. Un tipico kit potrebbe presentarsi così:



Vi occorreranno una serie di giunti ed altre soluzioni di natura plastica.



E se non vi siete affidati ad un kit preconfezionato, che sapete che funziona, dimensionato correttamente, una delle sfide più interes-

ti potrebbe essere proprio la progettazione di questi pezzi. Per i quali naturalmente vi servirà uno strumento CAD adeguato.

Il meccanismo di movimentazione degli assi è basato sull'utilizzo di una **cinghia di trasmissione**. Questo sistema propone alcune debolezze strutturali dovute sostanzialmente alla tenuta, nel corso del tempo, della cinghia che, sottoposta a sollecitazioni continue, soprattutto nel caso di materiali scadenti, può cedere, allentarsi e realizzare una movimentazione poco efficace.

Un sistema alternativo proposto da alcuni costruttori, che hanno realizzato stampanti più o meno professionali, prevede l'**uso della cremagliera**.



L'ipotesi che sorregge questa scelta è quella di una maggiore efficacia nel posizionamento, ed ovviamente nella più piccola dimensione ottenibile del pezzo. L'effetto, dunque, sarebbe relativo ad una più alta precisione realizzativa e così, infatti, viene giustificato anche un prezzo talvolta sensibilmente superiore rispetto alle soluzioni disponibili, soprattutto in ambito Maker.

Sapete che cos'è, di cosa stiamo parlando? Avete un'idea di come funziona?

Ma soprattutto: **potreste dire che si tratta di un meccanismo più o meno efficiente?**

I SERVO

Anche la scelta del motorino è fondamentale.

Qual è il passo che voglio che abbia? Quale l'impatto nella scelta del passo sulla resa di stampa? Come si pilota un motorino?

Questa volta facciamo qualcosa di diverso: stimoliamo i nostri lettori. Tanto per iniziare, perché il motorino funzioni perfettamente c'è bisogno che ci sia una libreria che lo gestisca e questa deve essere compatibile con il microcontrollore che stiamo scegliendo, i cui parametri comunque verranno specificati in un paragrafo dedicato per cui per il momento facciamo finta di poter ignorare la questione.

Vi suggeriamo un codice, NEMA17. Per molti sarà già familiare, per altri, invece non significherà niente. Se rientrate nella seconda categoria ecco cosa dovete fare per capire di che cosa stiamo parlando: fare una ricerca, studiare il manuale, identificare delle cifre di merito, confrontare i valori con qualcosa di simile, soprattutto guardare altri componenti per capirne le differenze e poi guardare cosa gli utenti in giro per la rete hanno fatto. Quando avrete finito, tornate qui con tutte le domande che vi sono venute in mente perché a quel punto sapete che cosa chiedere e sarà tutto più chiaro.

GLI UTENSILI

Per costruire la vostra stampante 3D è indispensabile che voi abbiate una serie di attrezzi che sono di uso comune e sempre presenti all'interno di un'officina meccanica o di un FabLab. Stiamo parlando, per esempio, di:

- **un set completo di cacciaviti;**
- **un set completo di cacciaviti di precisione;**
- **un set completo di brugole;**
- **un set completo di chiavi inglesi;**

- **calibro;**
- **pinze;**
- **pinze di precisione;**
- **livella.**

L'ELETTRONICA

Dal punto di vista elettronico, la stampante 3D è una macchina abbastanza semplice perché quello che vi occorre è una scheda basata su microcontrollore all'interno della quale sia possibile scrivere un firmware e che riesca a gestire tutto quello di cui abbiamo parlato fino a questo momento.

Motorini, ventole, display e tutto ciò che vorrete che faccia parte della vostra soluzione.

E allora questo vuol dire che dovrete fare particolare attenzione a una serie di parametri:

- **numero di GPIO disponibili per il microcontrollore;**
- **quantità di memoria disponibile;**
- **velocità di elaborazione dei dati;**
- **tensioni e correnti di alimentazione;**
- **tensioni e correnti gestite.**

E a proposito di questo, quando arrivate a progettare la parte elettronica dovete aver chiaro ciò di cui abbiamo parlato all'inizio ovvero tutte le caratteristiche che volete implementare perché sarà l'elettronica a gestirle e a consentirvi l'operatività di cui avete bisogno.

I MATERIALI

Questo paragrafo deve essere dedicato ai materiali per la stampa 3D... Quali sono? Che caratteristiche hanno? Che implica sulla lavorazione? Quale scelgo allora?

Prevalentemente la stampa 3D è divulgata come un metodo che lavora ed opera con materiale plastico e le soluzioni più comunemente conosciute solo il PLA, una plastica di origine

vegetale derivata dal mais, e l'ABS. La seconda è ancora una volta un materiale plastico, molto più resistente e durevole nel tempo, non sensibile all'acqua e all'umidità, e pertanto non è viene deformato né inficiato in termini di forma, colore e durata, ma risulta tossico.

Ecco per quale motivo suggeriamo all'inizio di tenere conto del fatto che i fumi prodotti dalla fusione del materiale potrebbero risultare tossici soprattutto in ambienti domestici di piccole dimensioni.

I materiali si propongono differenti anche per le temperature di lavorazione che nel primo caso si accostano in un intervallo compreso tra i **190 °C e 220 °C** mentre nel secondo caso vanno dai **210 °C fino addirittura a 280 °C**. Questi intervalli vanno comunque sottoposti a test e sono in effetti diversi, anche dei valori dichiarati, per ciascun produttore. Il vero test viene fatto durante la stampa con alcune impostazioni dedicate a questo.

La plastica, anche una volta realizzato il prototipo, però, resta dura. Non deformabile. Ecco per quale motivo esistono soluzioni che prevedono la stampa in gomma, delle quali una tra le più conosciute è proprio il NinjaFlex. Non ci sono limiti strutturali ma tenete presente che le temperature di estrusione sono diverse.

Altre possibilità riguardano l'utilizzo di materiali fluorescenti, in fibra di legno e così via dicendo. Al momento diversi sono i produttori che stanno studiando le più disparate ed innovative soluzioni in termini di materiali per cui stiamo tutti aspettando aggiornamenti a breve.

CONCLUSIONI

Forse sarete delusi alla fine di questo articolo perché non abbiamo effettivamente realizzato una stampante 3D con tutte queste informazio-

ni. Ma lo scopo non era questo. Lo scopo era quello di analizzare nel dettaglio **tutti i fattori** che **devono essere tenuti in considerazione** quando si sceglie di provare ad intraprendere questa strada. Anche il progettista più esperto, quello con anni di lavoro alle spalle sulla realizzazione di pezzi meccanici, incontrerebbe notevoli difficoltà nella creazione di una macchina a controllo numerico come questa. La precisione estrema con la quale è necessario trovare il compromesso e riprovare la progettazione dimostra che

“ la creazione di una stampante 3D è una sfida per VERI Maker.

Dal punto di vista elettronico tutto è "relativamente" più semplice, soprattutto se si è abbastanza abituati alla realizzazione di schede basate sul microcontrollori. Inoltre, in questo caso, non c'è da tenere conto di problematiche di compatibilità elettromagnetica... Ma siamo sicuri? Anche pilotare un motorino può comportare l'emissione di campi elettromagnetici, visto che certi valori di corrente non sono affatto trascurabili.

A proposito, abbiamo parlato dello spessore dei fili utilizzati per i collegamenti elettrici? **No. Ma dovremmo farlo.** Perché anche la sezione di un filo ed il valore massimo di corrente che può tollerare diventa un parametro di progetto tutt'altro che trascurabile.

E allora, quali sono le dimensioni più giuste secondo voi? Da cosa dipendono?

E ora che avete una visione più completa del problema, **come realizzereste la vostra stampante?**

L'autore è a disposizione nei commenti per eventuali approfondimenti sul tema dell'Articolo.
Di seguito il link per accedere direttamente all'articolo sul Blog e partecipare alla discussione:
<http://it.emcelettronica.com/cosa-serve-per-costruire-una-stampante-3d>

A questo numero di EOS-Book hanno partecipato:



Unendo la passione per la progettazione elettronica e le news da mondo tecnologico, nel Dicembre 2006 Emanuele dà vita a L'Elettronica Open Source, un blog condiviso rivolto alla fruibilità comune di risorse e progetti di elettronica e novità dal mondo dei semiconduttori e della tecnologia.

Chi accende una candela da te s'illumina anche lui, ma non ti lascia al buio - Share for life!

Emanuele
Founder
& Tech Supervisor



Al motto "Vola solo chi osa farlo", Piero è uno studioso appassionato ed uno smanettone curioso. Inizia come molti, smontando pc, si appassiona all'elettronica ed alla programmazione per poi oggi diventare progettista e maker con tantissimi modelli a cui tendere. Diverse le esperienze lavorative: dalle consulenze informatiche fino alla redazione di contenuti web passando per la realizzazione di progetti Open Source ed altro ancora.

Piero Boccadoro
Technical Writer



Da sempre appassionato di elettronica il percorso di studio non poteva che essere scontato: istituto tecnico indirizzo elettronica e laurea in ingegneria elettronica presso l'Università Politecnica delle Marche. Fortunatamente la passione si è trasformata in lavoro e attualmente mi occupo di sviluppo firmware per embedded. Molto spesso la curiosità mi spinge oltre gli ambiti lavorativi, motivo per cui sono venuto a conoscenza di Elettronica Open Source.

Luca Giuliadori
Contributor



Appassionato sin da piccolo per l'elettronica, la matematica ed il fai da te. E' programmatore di computer, insegnante di informatica e di matematica. Appassionato di numeri, è alla continua ricerca di grandi Numeri Primi. Ha scritto anche un libro sulla programmazione del PIC 16F84 con mikroBasic. E' titolare dell'azienda ElektroSoft, che si occupa di elettronica ed informatica. Si cura a tempo pieno di formazione ed insegnamento.

Giovanni Di Maria
Technical Writer



Bolognese, studi elettronici ma ha sempre fatto lo sviluppatore software a 360 gradi. Negli anni '80 e '90 scrive per diverse riviste cartacee come Elettronica Flash, Elettronica Giovane, Computer News, ecc. Dal 2000 residente a Fermo nelle Marche, ha aperto una software house di nome Brain & Bytes. Consulente e sviluppatore esterno di software per Trenitalia. Nel 2012 presenta un linguaggio di programmazione da lui sviluppato al Codemotion di Roma. Co-fondatore di Startupper Marche.

Gabriele Guizzardi
Contributor



Marco Giancola
Contributor

Laureato in Matematica e da sempre dotato di un'incommensurabile passione per la scienza e la tecnologia, ha iniziato la sua collaborazione con Elettronica Open Source nel settembre 2011 pubblicando un articolo di presentazione del suo e-book SATELLITI ARTIFICIALI. La sua successiva pubblicazione, Dimostrazioni matematiche della (non) esistenza di Dio, generò un tale clamore tra gli utenti del sito che furono postati più di 266 commenti.

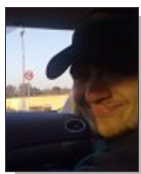
La libertà è poter affermare che due più due fa quattro. Se ciò è garantito, tutto il resto segue.

- George Orwell



XectX
Contributor

Studente della laurea magistrale in Ingegneria elettronica, appassionato di musica, attività fisica, scienza e tecnica. E' sempre impegnato in mille diverse attività ed affianca all'elettronica una modesta attività musicale. Le idee che gli passano per la testa sono sempre mille di più di quelle che riesce a realizzare.



Gransasso
Contributor

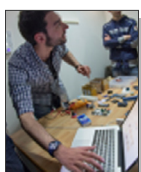
Laureato in ingegneria delle telecomunicazioni, interessi principali in elettromagnetismo, compatibilità elettromagnetica, elaborazioni numeriche, micro e nano elettronica.



slovati
Contributor

Stefano è laureato in Ingegneria Elettronica ed ha un'esperienza pluriennale nel settore dei sistemi embedded real-time, avendo partecipato a diversi progetti in campo avionico, automazione industriale, e telecomunicazioni. Ama lavorare direttamente sull'hardware e tenersi informato sulle nuove tecnologie elettroniche.

Il dubbio è l'inizio della sapienza (Cartesio)



Vincenzo Motta
Contributor

Ingegnere Elettronico, specializzando in Automation Engineering and Control of Complex Systems, ha sviluppato applicazioni per piattaforme modulari per la robotica presso STMicroelectronics. Da sempre affascinato dal mondo della robotica e dell'elettronica in generale, negli ultimi anni si è molto avvicinato al mondo dell'open source con Raspberry Pi e Arduino su cui ha tenuto workshop tra FabLab e Hackspace.

Non hai veramente capito qualcosa finché non sei in grado di spiegarlo a tua nonna
[Albert Einstein]