

ELEKTRONIKA PRAKTYCZNA

EP.com.pl

● Międzynarodowy magazyn elektroników konstruktorów ● listopad ● 11/2023 ●

Tylko Prenumeratorzy

- mają dostęp do artykułów przed ich publikacją w EP na www.ep.com.pl – **EP W TOKU**
- mają dostęp do materiałów dodatkowych, takich jak pliki źródłowe projektów na naszym serwerze **FTP** www.ulubionykiosk.pl/media

inspirujące, użyteczne projekty

- Dwukanałowy konwerter USB-C z układem FT2232H
- Uniwersalny przedwzmacniacz • Sygnalizator utraty ciągłości • Mikrowzmacniacz mocy 20 W na układzie PAM8320 • Moduł czterech wyjść HighSide dla RPi Pico
- Wskaźnik stanu emocjonalnego • MIDIXCV – konwerter MIDI do modułowego syntezatora analogowego
- multiLock

podzespoły, sprzęt, aplikacje

- Oscyloskop Siglent SDS1104X-U. Jeden z najlepszych budżetowych oscyloskopów cyfrowych • Regulator barwy dźwięku na bazie płytki ewaluacyjnej LPCxpresso55S28
- Aplikacje enkoderów obrotowych z zastosowaniem FPGA i układów SoC Microchip • Skazani na lit
- Laminat FR-4 – używany przez wielu, znany przez wielu • Płytki PCB zoptymalizowane do produkcji
- Wielowarstwowe PCB, od prostych do awangardowych
- Bezpieczne złącza do magazynów energii

tutoriale

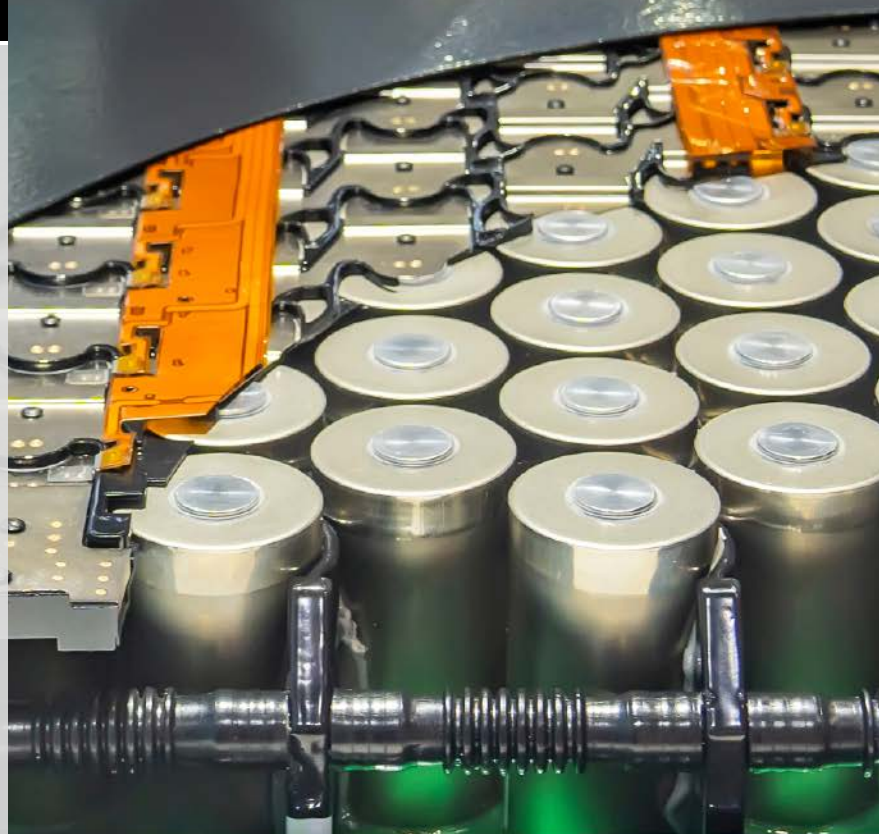
- Produkcja wielowarstwowych płytek PCB krok po kroku • GPSDO – jak GPS pilnuje czasu • Zagadnienia materiałowe w produkcji wielowarstwowych obwodów drukowanych • Kompleksowa produkcja elektroniki z zastosowaniem wielowarstwowych PCB • Wielowarstwowe płytki drukowane, co każdy projektant wiedzieć powinien? • Maksymalizuje sprawność elektroniki • Banki energii. Zarządzanie ciepłem i ochrona przed warunkami środowiskowymi

kursy

- Kurs FPGA Lattice. Wyświetlacz LCD multipleksowany

MAGAZYNOWANIE ENERGII ELEKTRYCZNEJ

TEMAT NUMERU



WIELOWARSTWOWE PŁYTKI PCB





**Zaprenumeruj
„Elektronikę Praktyczną”,
a zawsze dostaniesz
najnowszy numer wprost
do Twojej skrzynki!**

**na start
do 6* wydań gratis**

**po 5 latach
nieprzerwanej
prenumeraty
do 12* wydań gratis**

* Cena prenumeraty rocznej **na start** wynosi 207,90 zł. Przy zamówieniu prenumeraty dwuletniej za 340,20 zł oszczędność wynosi równowartość sześciu wydań „Elektroniki Praktycznej”.

Przedłużasz prenumeratę? Aby otrzymać zniżkę lojalnościową, przedłuż prenumeratę po zalogowaniu się do swojego panelu na www.ulubionykiosk.pl, gdzie znajdziesz atrakcyjną ofertę prenumeraty, która uwzględnia przysługujące Ci zniżki za lojalność. Po 5 latach nieprzerwanej prenumeraty otrzymasz **rabat 50%** na prenumeratę dwuletnią. Oferta dotyczy prenumeraty drukowanej.

**Wszystkie opcje prenumeraty i e-prenumeraty znajdziesz na stronie
www.UlubionyKiosk.pl**

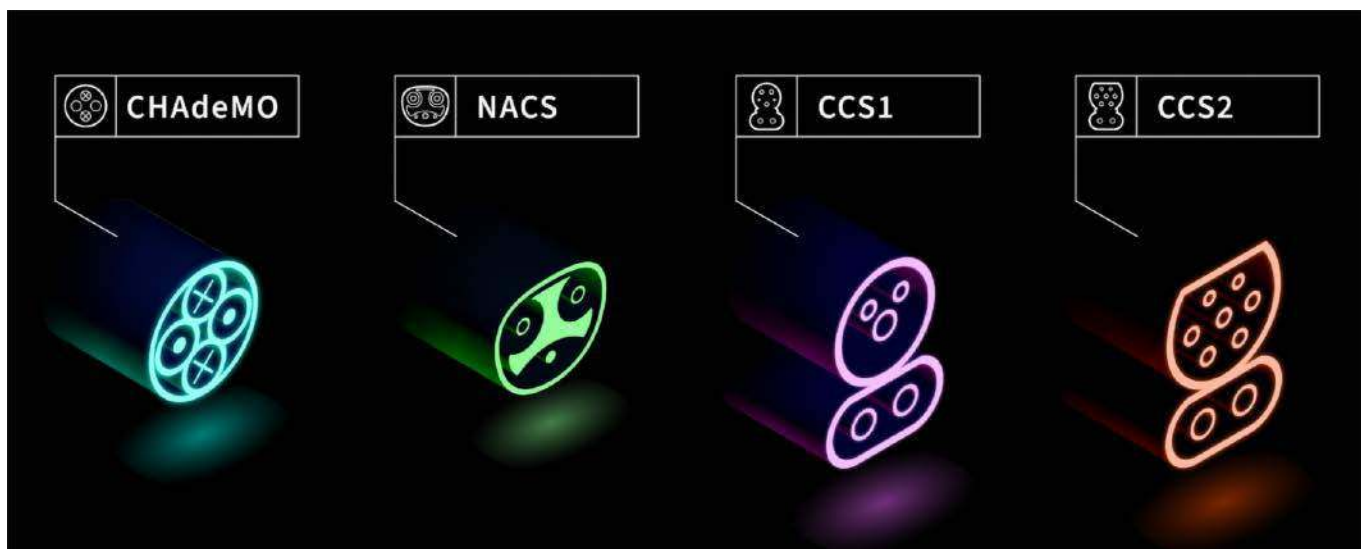
prenumerata@avt.pl

AVT-Korporacja sp. z o.o., ul. Leszczynowa 11, 03-197 Warszawa, konto 18 1050 1012 1000 0024 3173 1013

eprasa.pl a2e8019fb1

Megaherce i megawaty

Kryzys półprzewodnikowy mamy już za sobą. Od pewnego czasu producenci zaawansowanych układów scalonych pokrywają zapotrzebowanie na te komponenty generowane przez różne branże światowego przemysłu. Ponadto dostępne mikrokontrolery dobrze wpisują się w wymagania nowych aplikacji i są albo wystarczająco wydajne, albo odpowiednio wyspecjalizowane i pozwalają na realizowanie nowych, złożonych zadań (nawet jeśli oprogramowanie pozostawia wiele do życzenia w kwestii optymalizacji i wydajności). Oczywiście nie jest to wszechobecną regułą, ale nacisk na osiągnięcie coraz większej liczby megaherców wyraźnie osłabł.



Widoczny staje się nowy kierunek rozwoju, który dotyczy minimalizowania poboru mocy oraz maksymalizowania możliwości jej przetwarzania, dostarczania i magazynowania. Akumulatory zajmują szczególne miejsce w tym procesie – mają ogromny potencjał biznesowy, a zarazem są źródłem wielu kontrowersji; definiują możliwości wielu nowych rozwiązań, a jednocześnie obnażają słabości w zakresie standaryzacji infrastruktury energetycznej. Doskonale widać to na przykładzie branży motoryzacyjnej. Już teraz można zauważyć regionalną fragmentację standardów złączy ładowania (**rysunek powyżej**). W Europie przeważa system CCS1/CCS2, w USA jest NACS (znany jako standard Tesli), natomiast japońscy producenci opracowali złącza CHAdeMO.

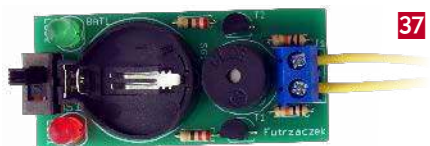
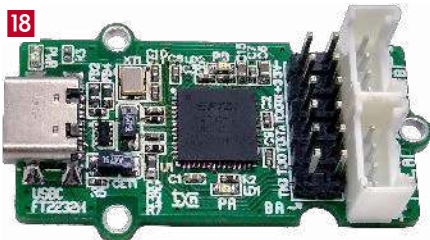
CCS1 nie sprawdził się jako uniwersalny standard. Koszty instalacji są wysokie (głównie ze względu na kable chłodzone cieczą, zaprojektowane na przesadną liczbę cykli), a wtyki są potężne i niewygodne w eksploatacji. CCS2 rozwiązuje część tych problemów, ale musi konkurować z rozwiązaniem NACS/Tesli, które jest obecnie znacznie bardziej praktyczne i łatwe w obsłudze. CCS mimo swoich wad jest otwartym standardem, spełniającym bardzo wyśrubowane normy (w tym bezpieczeństwa w komunikacji elektronicznej). NACS to ostatecznie rozwiązanie własnościowe Tesli i wiąże się z uzależnieniem od jednego producenta. To wszystko przypomina historię złączy ładowania do smartfonów – Lightning (Apple) oraz USB (wszyscy pozostali).



Zwiększenie pojemności akumulatorów najpewniej przyczyni się do uzyskania przewagi dla jednego ze standardów, ponieważ będzie wymagało ładowania większą mocą (aby zachować rozsądny czas ładowania). Standard CCS oferuje standardowo moc do 350 kW i jest dostosowany do architektury akumulatorów o napięciu 800 V (prąd do 500 A). W odpowiedzi na zapotrzebowanie na szybsze ładowanie wdrożono ładowarki CCS o mocy 400 kW i zademonstrowano ładowarki CCS o mocy 700 kW. NACS jest używane w samochodach Tesli od 2012 r. Konstrukcja ta została zaprojektowana dla architektury akumulatorów 400 V i jest znacznie mniejsza niż inne złącza szybkiego ładowania prądem stałym. Złącze NACS obecnie oferuje moc do 250 kW, jednak trwają prace nad wersją o mocy do 900 kW. Warto wiedzieć także, że dla branży elektrycznych samochodów ciężarowych, elektrycznych autobusów oraz elektrycznych pojazdów ciężkich wymagane jest nowe rozwiązanie w zakresie ładowania o dużej mocy. W rezultacie opracowywany jest obecnie Megawatt Charger System – MCS (**fotografia obok**). Pozwala na ładowanie z maksymalną mocą 3,75 MW (3000 amperów, przy napięciu 1250 V prądu stałego).

Płytki PCB mają wyjątkowo trudne zadanie – muszą sprostać niełatwym wymaganiom wynikającym zarówno z megawatowych mocy, jak i z megahercowych sygnałów. Dlatego warto poświęcić im więcej uwagi i zgłębić solidną porcję wiedzy, jaką przygotowaliśmy w tym wydaniu „Elektroniki Praktycznej”.

Damian Sosnowski



Nie przeocz

Nowe podzespoły	5
Dodaj do obserwowanych	12
Konkurs	28
Koktajl newsów	124

Projekty

Dwukanałowy konwerter USB-C z układem FT2232H	18
multiLock	21
Uniwersalny przedwzmacniacz (1)	29

Miniprojekty

Sygnalizator utraty ciągłości	37
Mikrowzmacniacz mocy 20 W na układzie PAM8320	38
Moduł czterech wyjść HighSide dla RPi Pico	40
Wskaźnik stanu emocjonalnego	42

Projekty SOFT

MIDIXCV – konwerter MIDI do modułowego syntezatora analogowego	111
--	-----

Moduły w aplikacjach

Regulator barwy dźwięku na bazie płytki ewaluacyjnej LPCxpresso55S28	108
--	-----

Temat numeru: Magazynowanie energii elektrycznej

Skazani na lit	50
----------------------	----

Prezentacje

Bezpieczne złącza do magazynów energii	46
Maksymalizuje sprawność elektroniki	49
Banki energii. Zarządzanie ciepłem i ochrona przed warunkami środowiskowymi	54
Wielowarstwowe PCB, od prostych do awangardowych	82
Płytki PCB zoptymalizowane do produkcji	84
Kompleksowa produkcja elektroniki z zastosowaniem wielowarstwowych PCB	86
Laminat FR-4 – używany przez wielu, znany przez nielicznych	90
Produkcja wielowarstwowych płytek PCB krok po kroku	94

Podzespoły

Aplikacje enkoderów obrotowych z zastosowaniem FPGA i układów SoC Microchip	56
---	----

Sprzęt

Oscyloskop Siglent SDS1104X-U. Jeden z najlepszych budżetowych oscyloskopów cyfrowych	58
---	----

Notatnik konstruktora

GPSDO – jak GPS pilnuje czasu	64
-------------------------------------	----

Elektronika w praktyce

Zagadnienia materiałowe w produkcji wielowarstwowych obwodów drukowanych	68
Wielowarstwowe płytki drukowane, co każdy projektant wiedzieć powinien?	98

Kursy

Kurs FPGA Lattice (13). Wyświetlacz LCD multipleksowany	116
Prenumerata	2
Od wydawcy	3
Hity następnego numeru	127

nowe podzespóły

Z kilkuset nowości wybraliśmy te, których nie wolno przeoczyć. Bieżące nowości można śledzić na www.elektronikaB2B.pl



100-woltowy tranzystor N-MOSFET do zastosowań w układach zasilania

Do oferty tranzystorów MOSFET firmy Toshiba, produkowanych w technologii U-MOS-X, wchodzi nowy model TPH3R10AQM, zaprojektowany do zastosowań w układach zasilania. Jest to tranzystor N-kanalowy o napięciu przebicia 100 V, charakteryzujący się szerokim obszarem bezpiecznej pracy (SOA) i małą rezystancją $R_{DS(ON)}$, wynoszącą maksymalnie 3,1 mΩ @ $V_{GS}=10$ V. Jest ona mniejsza o 16% od wcześniejszego modelu TPH3R70APL.

Tranzystor może pracować z maksymalnym prądem drenu 120 A. Charakteryzuje się małym ładunkiem QSW, pozwalającym na zastosowania w aplikacjach o dużej częstotliwości przełączania. Jego obszar SAO został poszerzony o 76% w porównaniu z tranzystorami wcześniejszej generacji, co jest istotne w obwodach hot-swap. Dodatkową zaletą jest mała powierzchnia montażowa, wynosząca 6,1×4,9 mm.

- Pozostałe parametry:
- QSW: typ. 32 nC,
- Qoss: typ. 88 nC,
- IDSS: maks. 10 μA @ $V_{DS}=100$ V,
- V_{th} : 2,5...3,5 V ($V_{DS}=10$ V, $I_D=0,5$ mA),
- V_{GSS} : ±20 V,
- PD: 210 W @ $T_c=25^\circ\text{C}$,
- EAS: 128 mJ.

www.toshiba.semicon-storage.com

Precyzyjny moduł nawigacyjny GNSS do robotów mobilnych

Firma u-blox prezentuje nowy, precyzyjny moduł nawigacyjny GNSS do robotów mobilnych, oznaczony symbolem NEO-F9P. Odbiera on równocześnie sygnały z satelitów konstelacji GPS, GLONASS, Galileo i BeiDou oraz obsługuje standard L1/L5 RTK, zapewniający centymetrową precyzję



pozycjonowania. Charakteryzuje się energooszczędną pracą (72 mA @ 3,0 V) i małymi gabarytami. Jest zamykany w obudowie o wymiarach 16×12×3,6 mm z wewnętrznym oscylatorem TCXO i pamięcią Flash.

NEO-F9P wymaga minimum elementów zewnętrznych. Może współpracować z antenami aktywnymi i pasywnymi, m.in. ANN-MB1. Zawiera interfejsy 2×UART, USB, SPI i DDC (kompatybilny z I²C). Jest przystosowany do pracy w zakresie temperatury otoczenia od -40 do +85°C.

www.u-blox.com



Wielostrefowy czujnik odległości z rodziny FlightSense z polem widzenia zbliżonym do kamery

VL53L7CX to wielostrefowy czujnik odległości z rodziny FlightSense o najszerszym polu widzenia spośród dostępnych obecnie odpowiedników. Może być stosowany m.in. do wykrywania obiektów i mapowania w automatyce domowej, sprzęcie AGD, robotach i fabrykach.

W przeciwieństwie do kamer, stosowanych niekiedy do podobnych zadań, czujniki ToF (*Time-of-Flight*) nie rejestrują obrazu, zapewniając tym samym pełną prywatność użytkowników. VL53L7CX pozwala znacznie rozszerzyć pole widzenia do wartości porównywalnej z tradycyjną kamerą, co zwiększa jego możliwości w zakresie wykrywania obecności i aktywowania urządzeń.

REKLAMA

BORNICO | Teraz większe MOŻLIWOŚCI

bornico.com.pl

- montaż kontraktowy elektroniki
- projektowanie urządzeń i systemów

Zakład Elektroniczny BORNICO

ul. Małczyńska 25
26-600 Radom
tel. +48 48 365 58 22
bornico@bornico.com.pl



Wielostrefowe czujniki FlightSense produkcji STMicroelectronics oferują duże możliwości w zakresie mapowania scen 3D i jednoczesnego pomiaru odległości do wielu obiektów, znajdujących się w różnych strefach. Możliwość pracy w wielu strefach wraz z funkcją sygnalizacji ruchu pozwalają na ich stosowanie np. do wykrywania obecności i śledzenia osób, ostrzegania o przekroczeniu strefy, zarządzania magazynem i miejscami parkingowymi itp. Z innych potencjalnych zastosowań należy wymienić korekcję zniekształceń trapezowych w projektorach i rozpoznawanie gestów przy użyciu pakietu oprogramowania STGesture.

VL53L7CX zawiera emiter VCSEL 940 nm. Pracuje z częstotliwością próbkowania 60 Hz. Umożliwia wykrywanie obiektów w maksymalnie 64 strefach (8×8), oferując w każdej z nich zasięg pomiaru od 2 do 350 cm. W trybie energooszczędnym pobiera zaledwie 5,4 mW mocy. Tryb autonomiczny umożliwia wyłączanie i budzenie mikrokontrolera host tylko po wykryciu ruchu lub po osiągnięciu określonej wartości progowej, np. minimalnej odległości do najbliższego obiektu.

VL53L7CX jest zamykany w obudowie LGA16 o wymiarach 6,4×3,0×1,6 mm, kompatybilnej pod względem rozkładu wyprowadzeń z czujnikiem poprzedniej generacji, VL53L5CX.

www.st.com



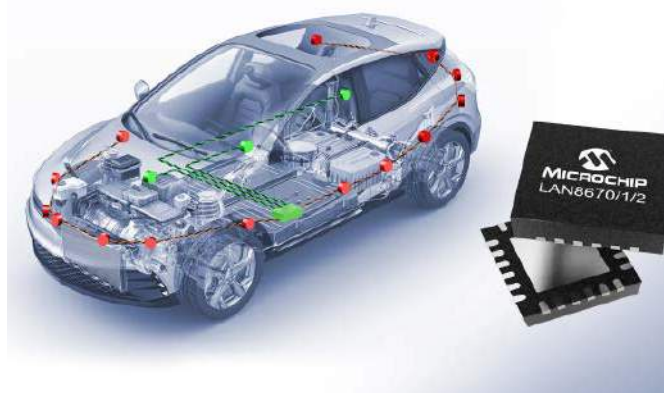
Samochodowa dioda LED RGB o dużej jasności i precyzyjnej kontroli barw

W ostatnich latach wprowadzono coraz więcej funkcji wspomagających kierowcę, np. automatycznej kontroli prędkości oraz wykrywania odległości do pojazdu i białej linii, co zwiększa zapotrzebowanie na chipowe diody LED RGB, zdolne do generowania wielu barw na tablicach przyrządów i zestawach wskaźników. Ponadto diody te są stosowane do oświetlenia dekoracyjnego w kabinie. Firma ROHM opracowała chipową diodę LED RGB o symbolu SMLVN6RGBFU, produkowaną na podłożu AlGaInP/InGaIn, zapewniającą dużą stabilność barw, wynikającą z precyzyjnej kontroli parametrów elementów R, G i B, takich jak długość fali i jasność. Uzyskała ona kwalifikację samochodową AEC-Q102. Może pracować w szerokim zakresie temperatury otoczenia od -40 do +100°C. Jest zamykana w obudowie SMD o powierzchni

	Barwa	Wartości maksymalne (+25°C)				Zakres temp. pracy	V _f @ 20 mA	Dominująca długość fali @ 20 mA	Intensywność @ 20 mA
		P _D	I _f	I _{TP}					
SMLVN6RGBFU	R		50 mA	100 mA			2,1 V	621 nm	750 mcd
	G	400 mW	40 mA	100 mA	-40...+100°C		3,3 V	525 nm	1800 mcd
	B		40 mA	100 mA			3,3 V	470 nm	430 mcd

3,5×2,8 mm i grubości zaledwie 0,6 mm. Wewnętrzne elementy R, G i B pracują na dominującej długości fali 621, 525 i 470 nm, a ich intensywność emisji wynosi odpowiednio 750, 1800 i 430 mcd.

www.rohm.com



Transceivery Ethernet 10BASE-T1S do zastosowań w motoryzacji

Microchip wprowadza do oferty swoje pierwsze transceivery 10BASE-T1S z kwalifikacją AEC-Q100 Grade 1, zaprojektowane do zastosowań w motoryzacji: LAN8670, LAN8671 i LAN8672. Mogą one pracować w temperaturze otoczenia od -40 do +125°C. Umożliwiają łączenie z samochodową siecią ethernetową urządzeń o małej szybkości transmisji, wymagających wcześniej własnych systemów komunikacyjnych. Zapewniają zgodność z wymogami normy ISO 26262 w zakresie bezpieczeństwa funkcjonalnego.

Możliwość podłączenia wielu transceiverów PHY Ethernet do wspólnej linii magistrali ułatwia wdrażanie aplikacji motoryzacyjnych na jednej dobrze znanej architekturze i oszczędza koszty dzięki uproszczeniu okablowania i zmniejszeniu liczby wymaganych portów w switchach. Transceivery serii LAN867x umożliwiają urządzeniom brzegowym korzystanie z Ethernetu i protokołu IP w celu łatwej komunikacji z resztą infrastruktury sieciowej. Oferują zaawansowane funkcje diagnostyczne, ułatwiające rozwiązywanie ewentualnych problemów z transmisją danych oraz funkcję uśpienia/budzenia, umożliwiającą korzystanie z trybów o zmniejszonym poborze mocy.

LAN8670, LAN8671 i LAN8672 różnią się rodzajem interfejsu (odpowiednio MII/RMII, RMII i MII). Pracują w trybie half-duplex z szybkością do 10 Mbps. Spełniają wymogi norm branżowych w zakresie kompatybilności elektromagnetycznej. Obsługują zestaw standardów TSN (Time-Sensitive Networking), pozwalając na równoczesny transfer danych w czasie rzeczywistym i obsługę aplikacji o dużej ilościach przetwarzanych danych (np. transmisji strumieniowej) przez wspólny przewód Ethernet, bez wzajemnych zakłóceń.

www.microchip.com

Pierwsze na rynku układy zarządzania energią do baterii zegarkowych

Nexperia prezentuje pierwsze na rynku układy zarządzania energią do baterii zegarkowych, pozwalające wydłużyć ich żywotność i zwiększyć energię wyjściową (power boost) przy pracy impulsowej.

Litowe baterie zegarkowe CR2032 i CR2025 ze względu na dużą gęstość energii i długą żywotność są powszechnie stosowane w aplikacjach o małym poborze mocy, w tym korzystających z komunikacji bezprzewodowej Wi-Fi, LoRa, Sigfox, ZigBee, LTE-M1 i NB-IoT. Jednak stosunkowo duża rezystancja wewnętrzna i mała szybkość zachodzenia reakcji chemicznych zmniejszają użyteczną pojemność baterii w warunkach obciążenia pulsacyjnego. Aby wyeliminować to ograniczenie, firma Nexperia opracowała układy zarządzania energią serii NBM7100 i NBM5100, wyposażone w dwa stopnie konwersji DC/DC oraz inteligentny algorytm uczenia się. W pierwszym etapie konwersji



energia jest powoli przesyłana z baterii do pojemnościowego elementu magazynującego. W drugim etapie zmagazynowana energia jest używana do wytwarzania impulsów, których napięcie może być programowane w zakresie od 1,8 do 3,6 V, a maksymalne natężenie prądu osiąga 200 mA. Inteligentny algorytm monitoruje energię zużywaną podczas powtarzających się cykli impulsów oddawanych do obciążenia i optymalizuje pracę pierwszego konwertera DC-DC, aby zminimalizować resztkowy ładunek w kondensatorze. Gdy nie są wykonywane cykle konwersji energii (stan czuwania), oba układy pobierają prąd o natężeniu mniejszym od 50 nA.

NBM7100 i NBM5100 są przeznaczone do pracy w temperaturze otoczenia od -40 do $+85^{\circ}\text{C}$, dzięki czemu nadają się do zastosowań komercyjnych i przemysłowych. Oferują funkcję ostrzegania systemu o zbliżającym się całkowitym rozładowaniu baterii. Ponadto funkcja ochrony przed spadkiem napięcia (*brownout protection*) wstrzymuje ładowanie kondensatora magazynującego, gdy akumulator zbliża się do końca okresu eksploatacji.

NBM7100 i NBM5100 pozwalają wydłużyć żywotność typowej litowej baterii zegarkowej nawet 10-krotnie przy jednoczesnym zwiększeniu nawet 25-krotnie szczytowego prądu wyjściowego. Pozwala to na zastosowania w wielu rodzajach urządzeń elektronicznych, w których dotychczas jedynym realnym źródłem zasilania były baterie AA lub AAA.

	Interfejs	Auto Start	Maks. napięcie kondensatora	Maks. prąd obciążenia
NBM7100A	I ² C	tak	11 V	200 mA
NBM5100A	I ² C	tak	5,5 V	150 mA
NBM7100B	SPI	nie	11 V	200 mA
NBM5100B	SPI	nie	5,5 V	150 mA

Układy serii NBMx100 występują w 4 wersjach różniących się rodzajem interfejsu (SPI lub I²C), opcjonalną funkcją autostartu oraz maksymalnym natężeniem prądu wyjściowego i maksymalnym napięciem kondensatora magazynującego. Ponadto warianty NBM5100A/B zawierają wyprowadzenie do równoważenia napięć, pozwalające na zastosowania w aplikacjach bazujących na superkondensatorach.

www.nexperia.com



Wzmacniacze niskoszumowe na pasmo 10 MHz...50 GHz zasilane napięciem 110/240 VAC

Fairview Microwave wprowadza do oferty nową serię niskoszumowych wzmacniaczy na zakres częstotliwości pracy 10 MHz...50 GHz, zasilanych napięciem 110/240 VAC. Są one produkowane na bazie półprzewodników GaAs, co zapewnia bardzo dobre właściwości szumowe i ułatwia wykrywanie słabych sygnałów. W zależności od wersji ich wzmacnienie wynosi od 25

	FMAM63022	FMAM63023	FMAM63024	FMAM63025
Pasmo	0,01...3 GHz	0,01...3,5 GHz	0,01...20 GHz	0,01...30 GHz
Wzmocnienie	36 dB	62 dB	29 dB	38 dB
P1dB	21 dBm	17 dBm	22 dBm	27 dBm
OIP3	36 dB	26 dB	29 dB	35 dB
NF	2,5 dB	1,3 dB	3,5 dB	3,7 dB
Złącze	SMA	SMA	SMA	2,92 mm

REKLAMA



Obudowa miniaturowa 1551W IP68

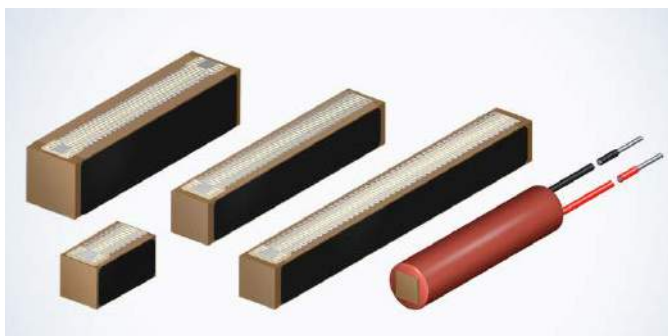
Dowiedz się więcej:
hammfg.com/1551w

eusales@hammfg.com • +44 1256 812812



do 60 dB. Wzmacniacze serii FMAM6302x są produkowane w aluminiowych obudowach klasy militarnej z wbudowanymi radiatorami, umożliwiającymi pracę w zakresie temperatury otoczenia od -40 do +85°C. Zapewniają odporność na udary do 20 g (11 ms) i wibracje do 25 g rms. W zależności od wersji zawierają złącza sygnałowe standardu SMA lub 2,92 mm.

www.fairviewmicrowave.com



Siłowniki piezoelektryczne do nanopozycjonowania i sterowania zaworami

Do oferty TDK wchodzi dwa nowe siłowniki piezoelektryczne, wykonane z cyrkonianu-tytanu ołowiu (PZT) z wewnętrzną elektrodą miedzianą, dostarczane jako pasywowane komponenty bez obudowy. COM30S5 (B58004M4030A020) i COM45S5 (B58004M4040A020) charakteryzują się bardzo szerokim zakresem dynamicznym, dużym stosunkiem siły do objętości oraz precyzją rzędu nanometrów. Uzyskano to dzięki opatentowanej technologii High Active Stack (HAS), która w porównaniu z innymi technologiami zapewnia dodatkowo dużą odporność na wilgoć i dłuższą żywotność.

Nowe siłowniki pracują z napięciem zasilania od -10 do +180 V, przy czym skok znamionowy jest osiągany przy napięciu +160 V. Zakres dopuszczalnej temperatury powierzchni wynosi od -40 do +160°C. Oba siłowniki zapewniają skok wynoszący odpowiednio 55 i 83 µm i charakteryzują się siłą 730 N. Ich wysokość wynosi odpowiednio 30 mm i 45 mm, a przekrój poprzeczny to 5,2×5,2 mm.

Seria	Kod zamówienia	Wymiary (L×W×H)	Skok @ +160 V	Siła obciążenia wstępnego	Siła blokowania
COM30S5	B58004M4030A020	5,2×5,2×30 mm	55 µm (±10%)	730 N	1400 N
COM45S5	B58004M4040A020	5,2×5,2×45 mm	83 µm (±10%)	730 N	1400 N
COM10S5	Z63000Z2910Z001Z78	5,2×5,2×10 mm	16 µm (±10%)	730 N	1400 N
COM27S3	Z63000Z2910Z001Z77	3,4×3,4×27 mm	47 µm (±10%)	320 N	600 N
COM30S7	Z63000Z2910Z001Z70	7,0×7,0×30 mm	55 µm (±10%)	1320 N	2600 N

TDK planuje wprowadzić na rynek jeszcze w 2023 roku trzy kolejne modele: COM10S5 (Z63000Z2910Z001Z78) o wysokości 10 mm i skoku 16 µm, COM27S3 (Z63000Z2910Z001Z77) o wysokości 27 mm i skoku 47 µm oraz COM30S7 (Z63000Z2910Z001Z70) o wysokości 30 mm i przekroju 7×7 mm, zapewniający skok 55 µm i siłę blokującą 2600 N.

Zakres zastosowań siłowników z nowej oferty obejmuje nanopozycjonowanie, sterowanie zaworami cieczy i gazów w technice procesowej i produkcję półprzewodników.

www.tdk-electronics.tdk.com

Czujnik wodoru do samochodowych systemów zarządzania akumulatorami

Posifa Technologies wprowadza do oferty czujnik wodoru PGS4100 do samochodowych systemów zarządzania akumulatorami (BMS), ułatwiający realizację obwodów zabezpieczających przed niekontrolowanym wzrostem temperatury (thermal runaway). Zapewnia on dokładny pomiar stężenia wodoru w powietrzu, mierząc zmianę przewodności cieplnej mieszaniny gazów. Umożliwia to skrócenie



czasu reakcji podczas uruchamiania alarmów awarii akumulatora w pojazdach elektrycznych (EV), zapewniając zgodność z branżowymi normami bezpieczeństwa.

PGS4100 zawiera czujnik wilgotności względnej i czujnik ciśnienia barometrycznego, pozwalające na równoważenie zmian przewodności cieplnej gazu, na którą wpływają zmiany wilgotności powietrza i wysokości. Obecnie jest produkowany w wersji z wyjściem analogowym 0,5...4,5 V i wyjściem cyfrowym I²C, a w przyszłości ma być też dostępny z interfejsami MODBUS/UART i CAN. Jest zamykany w obudowie o stopniu ochrony IP69K z kablem zakończonym wtykiem klasy motoryzacyjnej. Pracuje z napięciem zasilania 5,5 VDC, pobierając do 190 mW mocy.

Pozostałe parametry:

- zakres pomiarowy: 0...4% H₂,
- rozdzielczość: 10 ppm (wyjście analogowe); 2 ppm (wyjście cyfrowe),
- dokładność (0...20000 ppm H₂ @ 25°C): 1200 ppm,
- dokładność (>20000 ppm H₂ @ 25°C): 6% odczytu,
- dryft temperaturowy: maks. 1000 ppm (-40...+85°C),
- błąd długoterminowy: maks. 1200 ppm przez 4 lata,
- czas odpowiedzi: maks. 1 s (t₉₀),
- zakres wilgotności: 0...100 %RH,
- zakres ciśnienia roboczego: 30...120 kPa,
- zakres temperatury pracy: -40...+85°C.

www.posifatech.com

3-kanalowy sterownik silników BLDC, DC i szczotkowych o wyjściowym prądzie szczytowym 2 A

Do oferty firmy TDK wchodzi nowy sterownik silników bezszczotkowych (BLDC), szczotkowych (BDC) i krokowych, mogący pracować z maksymalnym prądem wyjściowym 2 A, przeznaczony do zastosowań w aplikacjach przemysłowych i motoryzacji. Zawiera on 3 wyjścia sterujące o rezystancji high-side i low-side wynoszącej typowo 0,7 Ω. Jego zastosowania obejmują głównie pompy, wentylatory i siłowniki małej mocy. HVC 5223C jest zamykany w obudowie QFN-24 o powierzchni 5×5 mm. Uzyskał kwalifikację AEC-Q101 Grade 1. Pod względem funkcjonalnym i rozkładu wyprowadzeń odpowiada wcześniejszej wersji 1-amperowej o symbolu HVC 5222C. Jego struktura obejmuje zestaw specjalizowanych bloków analogowych i cyfrowych, w tym 12-bitowy przetwornik A/C o cyklu 1 µs, komparatory fazy, wzmacniacz current-sense oraz transceivery UART i LIN z automatycznym adresowaniem metodą BSM.

Obecnie w ramach rodziny HVC dostępnych jest 7 sterowników, zawierających od 3 do 6 wyjść o prądzie szczytowym od 0,5 do 2 A. Wszystkie bazują na 32-bitowych jednostkach obliczeniowych ARM Cortex-M3 z 32 lub 64 kB pamięci Flash, 2 kB pamięci RAM, 512 B pamięci EEPROM i 256 B pamięci NVR

www.micronas.tdk.com



REKLAMA

COMPUTER CONTROLS

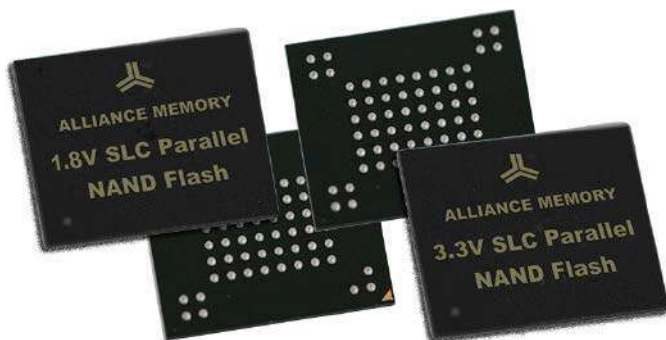
- / SPOTKANIA Z DOSTAWCAMI
- / SESJE TEMATYCZNE, PREZENTACJE, SEMINARIA
- / NETWORKING, DYSKUSJA I WYMIANA DOŚWIADCZEŃ

23 LISTOPADA

ZAPROJEKTUJ PRZYSZŁOŚĆ



WSPARCIE



Nowe pamięci NAND Flash SLC 1,8 V i 3,3 V z interfejsem równoległym

Aby sprostać wciąż dużemu zapotrzebowaniu na starsze równoległe pamięci NAND Flash SLC, firma Alliance Memory wprowadza na rynek nową serię układów o pojemności od 1 do 8 Gb, przeznaczonych na rynek motoryzacyjny, przemysłowy, komunikacyjny i elektroniki użytkowej. Pamięci AS9F z szyną I/O o szerokości $\times 8$ charakteryzują się krótkim czasem kasowania bloku, wynoszącym od 3 ms. Są zgodne ze specyfikacją ONFI 1.0. Podzielone na bloki z możliwością niezależnego kasowania, umożliwiają zachowanie ważnych danych, gdy stare dane są kasowane.

	Pojemność	Szyna	Vcc	Obudowa	Temp. pracy
AS9F31G08SA-25BIN	1 Gb	$\times 8$	3,3 V	FBGA-63	-40...+85°C
AS9F32G08SA-25BIN	2 Gb		3,3 V		
AS9F34G08SA-25BIN	4 Gb		3,3 V		
AS9F14G08SA-45BIN	4 Gb		1,8 V		
AS9F38G08SA-25BIN	8 Gb		3,3 V		
AS9F18G08SA-45BIN	8 Gb		1,8 V		-40...+105°C
AS9F14G08SA-45BAN	4 Gb		1,8 V		
AS9F18G08SA-45BAN	8 Gb		1,8 V		

Pamięci AS9F oferują niezawodność na poziomie 100 tys. cykli programowania/kasowania i 10 lat retencji danych. Obsługują mechanizm korekcji ECC. Są zamykane w obudowach FBGA-63 o wymiarach 11 $\times$ 9 $\times$ 1 mm. Występują w wersjach do zastosowań przemysłowych i motoryzacyjnych, charakteryzujących się zakresem temperatury pracy odpowiednio -40...+85°C i -40...+105°C. Dostępne są wersje o napięciu zasilania 1,8 V i 3,3 V.

www.alliancememory.com

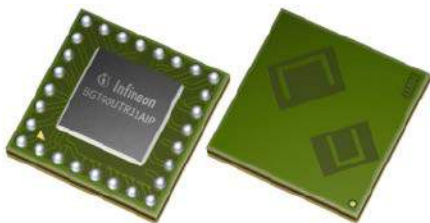
Zintegrowany czujnik radarowy do elektroniki użytkowej, aplikacji IoT i urządzeń medycznych

BGT60UTR11AIP

to czujnik radarowy rodziny Xsensiv o dużym stopniu integracji, produkowany w opracowanej przez Infineon technologii SiGe BiCMOS

B11, zapewniającej bardzo dobre parametry w.cz. Może on znaleźć zastosowanie w elektronice użytkowej, aplikacjach IoT i urządzeniach medycznych. Jest najmniejszym obecnie czujnikiem pracującym na częstotliwości 60 GHz, zamykanym w obudowie o wymiarach 4,05 $\times$ 4,05 $\times$ 0,86 mm z wewnętrzną anteną nadawczą i odbiorczą. Ze względu na niewielkie gabaryty jest polecany do urządzeń małego gabarytów, takich jak czujniki medyczne. Poza tym może znaleźć zastosowanie w elektronice konsumenckiej (laptopy, kamery, klimatyzatory, termostaty), robotyce i czujnikach poziomu. Pracuje w zakresie temperatury otoczenia od -20 do +70°C.

BGT60UTR11AIP umożliwia wykrywanie obecności i ruchu z milimetrową rozdzielczością w zakresie do 15 m. Komunikuje się



z mikroprocesorem przez interfejs API, służący jednocześnie do sterowania i transmisji danych. Zawiera wewnętrzne anteny o polu widzenia $\pm 60^\circ$, 12-bitowy przetwornik A/C o szybkości próbkowania do 4 MSps, czujniki mocy wyjściowej i temperatury oraz maszynę stanów, umożliwiającą akwizycję danych w czasie rzeczywistym, bez udziału mikroprocesora. Tryb broadcast umożliwia synchronizację pracy wielu układów.

www.infineon.com



Moduł nawigacyjny GNSS z chipem Teseo IV o dokładności pozycjonowania lepszej od 1 m

STMicroelectronics dodaje do oferty modułów nawigacyjnych GNSS nowy model Teseo-LIV4F, bazujący na chipie Teseo IV, obsługujący konstelacje satelitów GPS, Galileo, Glonass, BeiDou i QZSS. Zapewnia on dokładność pozycjonowania lepszą od 1 m. Jest dostarczany wraz z oprogramowaniem układowym Teseo-LIV4FSW, dzięki czemu użytkownicy mogą korzystać z aplikacji Teseo Suite do bezpłatnej konfiguracji i aktualizacji oprogramowania układowego. Sprawdzone środowisko projektowe skraca czas projektowania, a mała objętość modułu i niski koszt sprawiają, że Teseo-LIV4F jest idealny do szerokiej gamy zastosowań, w tym śledzenia towarów, systemów antykradzieżowych i pobierania opłat, lokalizacji ludzi i zwierząt, śledzenia pojazdów, połączeń alarmowych, zarządzania flotą pojazdów, diagnostyki i transportu publicznego. Kolejne zalety to małe gabaryty i energooszczędna praca.

Teseo-LIV4F charakteryzuje się czułością -162 dBm w trybie śledzenia, co zapewnia dużą dokładność pozycjonowania również przy słabych sygnałach z satelitów. Zintegrowany oscylator TCXO oraz dodatkowy oscylator do taktowania zegara RTC pozwalają zapewnić dużą dokładność pozycjonowania i krótki czas pierwszej akwizycji (TTFF).

Układ jest zamykany w obudowie LCC-18 o powierzchni 10,1 $\times$ 9,7 mm. Pracuje z napięciem zasilania od 3,0 do 3,63 V, pobierając zaledwie 10 μ A prądu w trybie standby i 48,8 mA w trybie śledzenia w paśmie L1/L5. Jest przystosowany do pracy w przemysłowym zakresie temperatury otoczenia od -40 do +85°C, co pozwala na zastosowania również na zewnątrz budynków oraz w warunkach przemysłowych.

www.st.com

Przemysłowe monitory dotykowe General Touch – idealne rozwiązanie do infokiosków

Monitory przemysłowe firmy General Touch to połączenie zaawansowanych technologii i innowacyjnego designu. Od dwóch dekad firma oferuje produkty, które sprawdzają się w miejscach, gdzie kluczowe są niezawodność i wytrzymałość – takich jak linie produkcyjne, systemy zarządzania procesami przemysłowymi, a także w miejscach publicznych, takich jak infokioski czy punkty sprzedaży. Warto przy tym dodać, że produkty General Touch mają certyfikaty CE i występują w wersjach





z wbudowanym komputerem i interfejsem HDMI, co sprawia, że jest to rozwiązanie łatwe w implementacji.

Monitor przemysłowy to wyświetlacz LCD-TFT umieszczony w szczelnej obudowie, który charakteryzuje się wyższymi parametrami, trwałością i bezawaryjnością, a także łatwością obsługi, co znacznie ułatwia pracę operatorów. Jednym z kluczowych parametrów monitorów przemysłowych General Touch jest wysoka rozdzielczość, która zapewnia doskonałą jakość obrazu. Monitory te oferują również pełny zakres kolorów oraz szerokie kąty widzenia, co pozytywnie wpływa na czytelność prezentowanych informacji.

Ponadto monitory General Touch są znane ze swojej wytrzymałości. Zostały zaprojektowane tak, aby wytrzymać trudne warunki

przemysłowe, jak wilgoć czy wibracje. Ich solidna konstrukcja gwarantuje długą żywotność i niezawodność. Ich technologia dotykowa wykorzystuje przede wszystkim rozwiązania bazujące na technologii pojemnościowej, co pozwala na precyzyjną i niezawodną detekcję dotyku, nawet w trudnych warunkach (np. w rękawicach roboczych).

Oprócz tego monitory przemysłowe General Touch wyróżniają się niskim poborem energii (18...22 i 28 W), co jest niezwykle istotne w kontekście długotrwałego użytkowania w przemyśle. Częściowo na niski pobór mocy wpływa zastosowanie nowoczesnych sposobów kontroli podświetlenia ekranu, czyli rozwiązań *local dimming*.

W tabeli zestawiono najważniejsze parametry wybranych modeli General Touch z oferty Unisystemu.

parametr	model OTL195-RPC05-UHPCD	model OTL175-RC6000-UH5100
przekątna	19"	17"
rozdzielczość	1280×1024 px	1280×1024 px
obszar aktywny	376,32×301,06 mm	341,92×274,32 mm
jasność	210 cd/m ²	415 cd/m ²
kontrast	1000:1	1000:1
kąty obserwacji	85/85/80/80°	80/80/75/75°
interfejs	DP, DVI, HDMI, VGA, USB-B	HDMI, VGA, USB-B
wymiary modułu	430,0×353,0×37,6 mm	380,00×317,50×59,60 mm
podświetlenie	LED	LED
proporcje obrazu	5:4	5:4
typ obudowy	closed frame	open frame
zakres temperatur pracy	0...40°C	0...50°C
panel dotykowy	tak, pojemnościowy	tak, pojemnościowy

www.unisystem.com/pl

REKLAMA

Nie przegap listopadowego wydania „Elektroniki dla Wszystkich”



przejrysz i kupisz na
www.ulubionykiosk.pl

dodaj do obserwowanych

Przedstawiamy redakcyjny wybór najciekawszych projektów spośród ostatnio anonsowanych w internecie. Są to projekty na różnych etapach realizacji. Warto się zapoznać z projektami zakończonymi i śledzić realizację projektów niegotowych, by czerpać z nich inspirację do własnych prac.

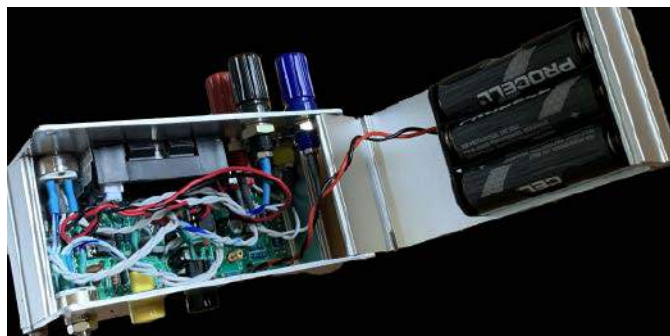


Minimalistyczny miliomierz

Oto prosty, przenośny miliomierz. Autor, projektując to urządzenie, przyjął następujące założenia:

- układ musi być w pełni zintegrowany – bez konieczności podłączania dodatkowych elementów,
- miliomierz ma korzystać z niedrogich i łatwych w zakupie elementów,
- zasilanie ma pochodzić z trzech baterii AA,
- układ realizuje tylko tryb pomiaru czteroprzewodowego (Kelvina),
- prąd pomiaru 100 mA,
- stały zakres 1 Ω z możliwością przekroczenia o 20%,
- dokładność ± 1 m Ω w zakresie temperatur 15...35°C,
- 4,5-cyfrowy wyświetlacz dla rozdzielczości 0,1 m Ω i częstotliwości odświeżania 3 Hz,
- analogowe niez izolowane wyjście 1 V/ Ω ,
- wskaźnik stanu baterii z trzema poziomami,
- zakres napięcia zasilania od 3 V do 5 V.

Urządzenie jest przeznaczone do użytku na stole lub jako przenośne, jako stosunkowo uniwersalny instrument pomiarowy, pomimo swojej prostoty i niskiego kosztu. Podstawowym celem projektu jest



wykazanie, że bardzo podstawowy układ może dobrze sprawować się w praktyce. Jest to również krok eksploracyjny w kierunku przyszłych ulepszeń. Autor obecnie pracuje nad lepszą i bardziej dopracowaną wersją.

Ten prosty miliomierz zaprojektowano z zasilaniem baterijnym w zamyśle. Dodatkowo, autor zaplanował, że układ ma bazować na przystępnych cenowo powszechnie dostępnych podzespołach. Jednakże zachowano pewną elastyczność: droższe substituty tych podzespołów miałyby mieć lepsze działanie. Kluczową rolę w projekcie odgrywają wzmacniacze operacyjne i komparatory, które realizują większość jego funkcji.

Projekt został zainspirowany projektem miernika firmy Electrolab, który został zaprezentowany na YouTube. Całkowity koszt tego urządzenia to około 15 dolarów.

To tylko prototyp! Projekt ten ma jeszcze pewne problemy, które nie są obce tym, którzy po raz pierwszy składają nowy projekt. Błędem było zaprojektowanie płytki PCB na podstawie wymiarów obudowy z aukcji na eBay – PCB jest zbyt szeroka (o 1...2 mm) i musiała być wyrównana, aby zmieściła się między przednią a tylną płytą obudowy.

<https://hackaday.io/project/191204-sdo-milliohm-meter-v1>



Miniaturowy miernik panelowy z wyświetlaczem OLED

Zaprezentowany projekt to programowalny miernik panelowy na bazie Arduino z wyświetlaczem OLED. Zawiera układ monitorowania zasilania INA226 do precyzyjnego pomiaru napięcia i prądu i prezentuje wyniki na ekranie z kontrolerem SSD1306. W odróżnieniu od innych tanich mierników panelowych, które często są bardzo niedokładne lub mają potencjometr wpływający na wyniki odczytu, ten miernik panelowy został zaprojektowany tak, aby dostarczyć dokładne pomiary zarówno napięcia, jak i prądu. Jest to małe, łatwe w użyciu urządzenie, które pozwala na pomiar mocy oraz zmianę jednostek dowolnej wielkości.

Omawiany miernik panelowy bazuje na układzie scalonym INA226, który umożliwia pomiar napięcia do 36 V i prądu do 5 A za pomocą

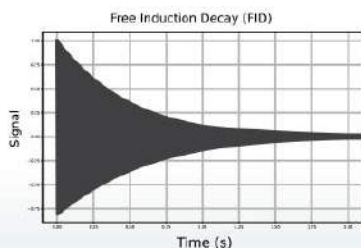


rezystora 5 mΩ po stronie wysokiej linii zasilania. Dzięki temu układowi odczyty są bardzo dokładne, z błędem wzmocnienia wynoszącym jedynie 0,1%. Stabilizator LDO 3,3 V zasila całą płytę, w tym wyświetlacz OLED. Stabilizator ten działa w zakresie napięcia wejściowego od 5 V do 40 V.

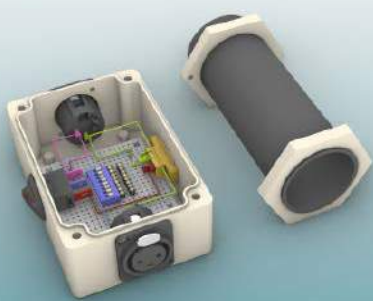
Płytkę miernika zawiera układ ATmega328P, który został zaprogramowany do pomiaru i wyświetlania napięcia od 0 V do 36 V i prądu od 0 A do 5 A. Precyzja pomiarów napięć pozwala mierzyć napięcie od 5 mV a prądy od 5 mA. Mimo zaawansowanych funkcji miernik panelowy cechuje się niewielkimi rozmiarami – zaledwie 27×28 mm, co sprawia, że jest łatwy w montażu w dowolnym urządzeniu. Ponadto w mierniku znajduje się stabilizator napięcia 3,3 V, z opcją jego obejścia, co pozwala na elastyczne zarządzanie zasilaniem układu.

W celu dostosowania działania płytki do naszych potrzeb czy chęci zmiany parametrów pomiarowych użytkownik może użyć programatora z konwerterem FTDI i po prostu podłączyć go do pinów wejściowych płytki.

<https://hackaday.io/project/191203-mini-oled-panel-meter>

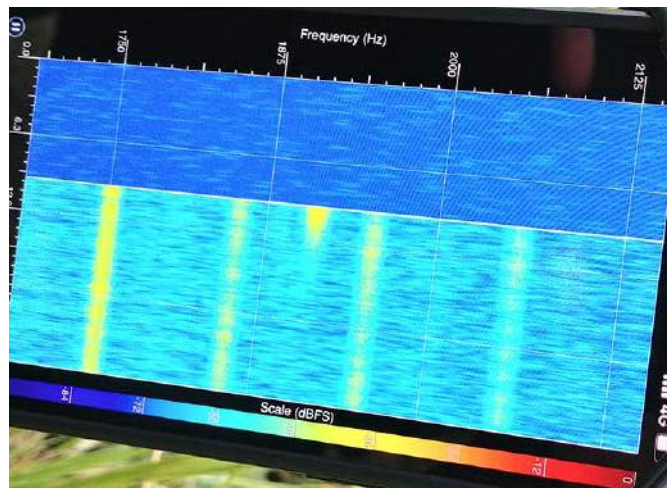


Nuclear Magnetic Resonance for Everyone



Jądrowy rezonans magnetyczny dla każdego

Jądrowy rezonans magnetyczny (NMR) to potężna technika leżąca u podstaw badań chemicznych i fizycznych. Oferuje cenny wgląd w struktury atomowe i molekularne poprzez badanie właściwości magnetycznych jąder atomowych. Projekt ten ma przeprowadzić przez proces budowy przyjaznego dla użytkownika, przenośnego urządzenia, które może niezawodnie mierzyć sygnały NMR w ziemskim polu magnetycznym.



Niezbędnie wiadomo, że wykrywanie podstawowego sygnału rezonansu magnetycznego (MRI) jest stosunkowo prostym zadaniem. W rezonansie magnetycznym częstotliwość Larmora ma kluczowe znaczenie dla zrozumienia, w jaki sposób momenty magnetyczne spinów jądrowych wyrównują się zgodnie lub przeciwnie do zewnętrznego pola magnetycznego oraz w jaki sposób następnie poruszają się one wokół kierunku pola magnetycznego. Ta precesja stanowi podstawę do generowania sygnałów stosowanych do tworzenia obrazów MRI. Podobnie częstotliwość Larmora jest niezbędna do zrozumienia spektroskopii magnetycznego rezonansu jądrowego (NMR), potężnej techniki analitycznej, która wykorzystuje właściwości magnetyczne jąder atomowych do badania struktury, dynamiki i interakcji cząsteczek. W NMR częstotliwość Larmora określa szybkość precesji jąder wokół kierunku pola magnetycznego, a precesja ta służy do generowania sygnałów stanowiących podstawę widma NMR.

Podobnie jak w MRI i spektrometrach NMR, które zwykle zawierają nadprzewodzące, potężne magnesy wykonane z materiałów ziem

REKLAMA

ZAJRZYSZ NA TE STRONY

All In One

- Projektowanie i wykonywanie
 - modeli kaskasów i obudów na drukarce 3D
 - transformatorów i induktorów
 - prototypów PCB
- Modelowanie 3D modułów i urządzeń
- Projektowanie urządzeń zasilających



Feryster - producent elementów EMC

FERYSTER

www.feryster.pl

RACK i Eurocarta 19" Wyposażenie szaf 19"

www.obudowa.pl

Producent obudów dla elektroniki tel. 032-230-2301



www.piekarz.pl

części elektroniczne
sprzedaz@piekarz.pl tel. 22 599 49 70



www.gamma.pl

PODZESPOŁY ELEKTRONICZNE



rzadkich, niektóre atomy w cieczach również naturalnie podlegają precesji w ziemskim polu magnetycznym. Częstotliwość Larmora zależy od siły przyłożonego pola, więc atomy wodoru w polu ziemskim wykazują bardzo niską częstotliwość precesji. Jest to bardzo dogodny zbieg okoliczności, gdyż ta częstotliwość mieści się w zakresie odpowiednim dla standardowego wzmacniacza audio.

Autor opisuje w artykule, jak zastosować łatwo dostępny interfejs audio na USB, aby ominąć ogromną złożoność zwykle napotykaną podczas projektowania niezawodnego, przenośnego wzmacniacza o bardzo niskim poziomie szumów i przetwornika ADC odpowiedniego do pomiaru sygnałów z atomów wodoru w ziemskim polu magnetycznym. Używając ręcznie obsługiwanego przełącznika trójdrożnego, omija się złożoność programowania impulsów NMR i eliminuje potrzebę stosowania cyfrowych obwodów sterujących.

Oprócz konstrukcji wzmacniacza i przetwornika ADC, konstrukcja cewki polaryzacyjnej i detekcyjnej jest często kolejnym kłopotliwym szczegółem technicznym w projektowaniu systemu NMR korzystającym z pola magnetycznego Ziemi (EFNMR). Wiele projektów wymaga dwóch lub więcej cewek i cienkiego drutu z tysiącami zwojów. Jednak ta konstrukcja zawiera krótką dwuwarstwową cewkę złożoną ze stosunkowo łatwego do nawinięcia drutu o średnicy 0,55 mm, dzięki czemu projekt jest bardziej dostępny.

Do wizualizacji odbieranych sygnałów służy popularny interfejs audio z wieloplatformowymi sterownikami, który można podłączyć bezpośrednio do telefonu, tabletu czy laptopa. Dane są analizowane bezpłatnymi narzędziami do analizy widmowej z wyjściami kaskadowymi o wysokiej wydajności (STFT). Tego typu narzędzia idealnie nadają się do pomiarów na żywo, nawet w pomieszczeniach w środowiskach miejskich o średnim poziomie hałasu, bez konieczności stosowania dodatkowego filtrowania sprzętowego lub ekranowania. Większość języków programowania ma biblioteki dla urządzeń audio USB, dzięki czemu projekt nadaje się do tworzenia niestandardowego oprogramowania.

Ponieważ EFNMR obejmuje pomiar niezwykle subtelnych ruchów magnetycznych atomów, interferencje mają kluczowe znaczenie dla sukcesu. Różnica między pomyślnym wykryciem sygnału a jego niepowodzeniem może sprowadzać się do przemieszczenia cewek o zaledwie 30 cm. Kluczowe znaczenie ma możliwość swobodnego obracania cewki w celu poruszania się w lokalnym polu magnetycznym. Cewka detekcyjna EFNMR musi być wysoce mobilna, aby zapewnić użyteczność systemu. Opisany system zawiera powszechnie spotykany, niedrogi aluminiowy statyw do aparatu z niewielkimi modyfikacjami, tworząc stabilną i lekką platformę pomiarową. Ta przenośność sprawia, że system doskonale nadaje się do pomiarów w terenie, zwiększając praktyczność EFNMR w zastosowaniach takich jak monitorowanie środowiska, nauki geograficzne, nauki o żywności i eksploracja zasobów.

<https://shorturl.at/JNOQZ>



Inteligentna przetwornica DC/DC 12 V typu Buck

Zaprezentowany projekt to przetwornica, która redukuje napięcie z baterii litowo-jonowych (od 16 V do 22 V DC) do stabilizowanego 12 V. Głównym elementem tej płytki jest scalony stabilizator napięcia firmy Texas Instruments – LM2596, wspomagany przez układ INA219 do dokładnego monitorowania prądu układu. Początkowo projekt został opracowany w celu zasilania kamery astrofotograficznej ZWO, jednak możliwości tej przetwornicy sięgają daleko poza to zastosowanie. Układ może być używany do różnych aplikacji, które wymagają stabilnego źródła napięcia 12 V DC o maksymalnym prądzie do 3 A, umożliwiając jednocześnie monitorowanie wejścia i wyjścia zasilania.

Obecnie projekt sprzętu przeszedł drugą rewizję, uwzględniając kluczowe poprawki po zidentyfikowaniu początkowych błędów. Aby uzyskać pełne informacje oraz dostęp do szczegółów projektu i plików źródłowych, należy odwiedzić stronę projektu na GitHubie.

<https://shorturl.at/suy03>

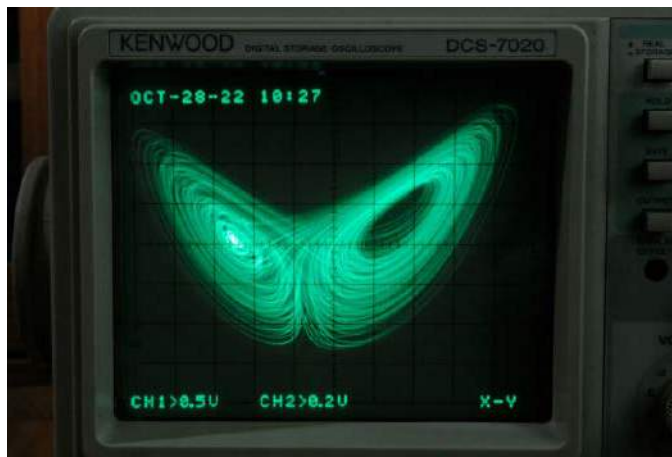


Robot pilnujący, który śledzi ludzi i psy

W okresie poprzedzającym wydanie nowej gry z serii Zelda, autor tej ciekawej konstrukcji, zdał sobie sprawę, że można zbudować stacjonarnego robota-strażnika z serwo mechanizmem i kamerą. Dodając trochę uczenia maszynowego, udało mu się prawić, by strażnik wykrywał przedmioty, ludzi lub zwierzęta domowe i podążał za nimi, obracając głowę.

Ten projekt pozwala opracować funkcjonalnego opiekuna z serwo mechanizmem, kamerą, kilkoma diodami LED oraz usługą Viam ML Model i Vision Service do rozpoznawania obrazu.

<https://shorturl.at/wEGR0>



Analogowy komputer rozwiązujący układ równań Lorenza

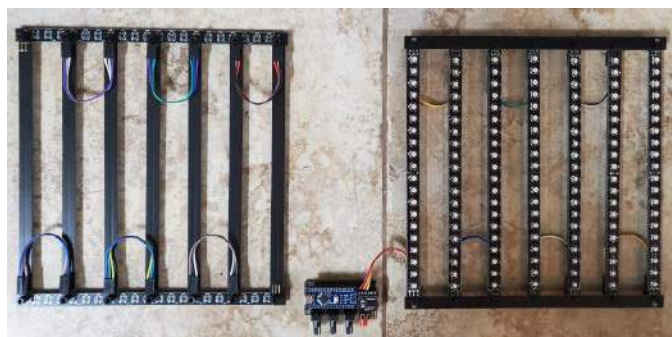
Czym w ogóle jest komputer analogowy? Jednym z głównych zadań obwodów analogowych jest rozwiązywanie problemów matematycznych, takich jak rozwiązywanie nieliniowych równań różniczkowych czy analizowanie ich charakterystyki poprzez analizę płaszczyzny fazowej napięcia wyjściowego za pomocą oscyloskopu lub plotera analogowego. Przykładem takiego problemu jest słynne nieliniowe równanie różniczkowe, tak zwany atraktor Lorenza. W ramach tego projektu autor pokazuje, w jaki sposób możliwe jest rozwiązanie tego problemu za pomocą obwodu analogowego.

Atraktor Lorenza to układ równań opisujących dynamiczne zachowanie atmosfery, który pokazuje chaotyczne zjawiska zawarte w zmianach meteorologicznych i jest znany, jako „efekt motyla”. Atraktor Lorenza składa się z trzech nieliniowych równań różniczkowych. Mają one trzy parametry: sigma, b i r, które określają charakterystykę systemu. Powszechnie przyjmowane wartości to sigma=10, r=28 i b=8/3.

Przed zbudowaniem obwodu analogowego autor chciał go symulować, aby zrozumieć zachowanie systemu. W tym celu użył oprogramowania Octave z funkcją atraktora Lorentza zadeklarowaną w lorenz.m, gdzie x jest trójwymiarowym zbiorem wektorów, a x(1), x(2) i x(3), które reprezentują, odpowiednio, x, y i z w oryginalnym zestawie równań.

Jak ten zestaw równań można „przetłumaczyć” na obwód analogowy? Jest to dokładnie zobrazowane w diagramie pokazanym w projekcie. W układach analogowych najwygodniejszymi układami do wykonywania obliczeń są układy zaprojektowane na wzmacniaczach odwracających, które są prostsze w formie niż wzmacniacze nieodwracające, a także unikają zniekształceń i szumów współbieżnych, co jest szczególnie ważne, gdy stosujemy wzmacniacze operacyjne o kiepskich parametrach szumowych.

<https://shorturl.at/cjwGS>



Siedmiopasmowy analizator widma dla sygnałów audio

Zaprezentowany system to siedmiopasmowy spektrometr, przeznaczony do analizowania sygnałów audio, to znaczy o częstotliwościach



w zakresie od 20 Hz do 20 kHz. Pozwala on na wizualizowanie widma sygnału. Widmo to reprezentacja częstotliwościowa sygnału, która pokazuje, jakie częstotliwości są obecne w danym sygnale i jak mocno są one wyrażone.

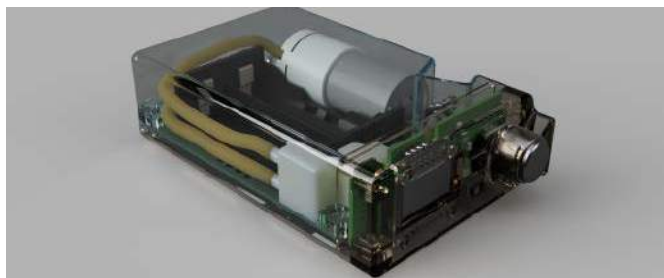
Konstrukcja zawiera układ scalony MSGEQ7, który jest scalonym analizatorem widma dla sygnałów audio. Ten analizator widma ma siedem pasm częstotliwościowych, które umożliwiają podzielenie widma na tyleż zakresów. W każdym z tych pasm można odczytać intensywność sygnału dla poszczególnych częstotliwości. Dzięki temu użytkownik może łatwo zidentyfikować dominujące częstotliwości w analizowanym sygnale i ocenić charakterystykę częstotliwościową dźwięku.

Za sterowanie układem MSGEQ7 odpowiedzialny jest moduł Arduino Nano. Wyświetla on też odpowiednie pasma na paskach LED WS2812B. Układ MSGEQ ma wbudowany generator zegarowy, więc nie ma potrzeby stosowania zewnętrznego oscylatora kwarcowego. Jest to idealny sposób na rozpoczęcia nauki i zabawy z funkcjami FastLED i NeoMatrix oraz pozwala zrozumieć, jak działają analizatory widma.

Jedynym ograniczeniem Arduino Nano jest mała pamięć Flash mikrokontrolera – jedynie 32 kB, która ledwo wystarcza na załadowanie kodu tego programu. Szkic dołączony do tego projektu ma jednak 16 różnych wzorów i kolorów dla FastLED, które przełącza się za pomocą przycisku lub pozwala się im automatycznie zmieniać co 10 sekund na następny wzór.

Projekt jest na tyle prosty, że można go przetestować na małej płytce prototypowej. Autor zaprojektował też płytki drukowane specjalnie dla tego układu, które pokazuje na stronie projektu.

<https://shorturl.at/efIV1>



Auto-Inflate – sterownik do kamizelki pneumatycznej do stymulacji sensorycznej

Urządzenie zostało zaprojektowane do sterowania nadmuchiwanymi kamizelkami sensorycznymi lub antystresowymi, takimi jak kamizelka The Squease Vest. Oczywiście, to urządzenie może działać także z innymi kamizelkami tego rodzaju, jeśli takie są dostępne. Auto-Inflate ma być noszone przez osobę i ma stosunkowo niewielkie wymiary, dzięki czemu można je ukryć w kieszonce lub torbie.

Niestety na ten moment firma Squease poinformowała, że zamierza zakończyć działalność. To urządzenie zostało specjalnie



zaprojektowane do współpracy z kamizelkami tej firmy, ale autor stwierdza, że system powinien działać z innymi kamizelkami tego rodzaju.

Prezentowana konstrukcja to udoskonalona wersja poprzedniego projektu Automatic Timed Inflation Device, który nigdy nie został sfinalizowany. Rozbudowa tej konstrukcji doprowadziła do stworzenia wielofunkcyjnej platformy, której główną funkcją jest nadmuchiwanie kamizelki sensorycznej. Najnowsza wersja eliminuje interfejs dotykowy na rzecz prostszego interfejsu z wyświetlaczem OLED i enkoderem. Nadmuch i spuszczenie powietrza są monitorowane za pomocą wbudowanego czujnika ciśnienia. Jedną ważną różnicą będzie możliwość zastosowania interfejsu haptycznego i systemu sprzężenia zwrotnego. Pozwoli to na aktywację za pomocą prostych dotknięć, aby rozpocząć cykl nadmuchu.

Auto-Inflate będzie miał bezprzewodową kontrolę za pośrednictwem Bluetooth lub Wi-Fi, co zwiększy użyteczność dla osób niepełnosprawnych, które korzystają z ubrań nadmuchiwalnych. Istotnym dodatkiem jest przycisk awaryjnego wyłącznika, który odłącza zasilanie komponentów wewnętrznych. Zapewnia to ręczne odcięcie zasilania w przypadku potencjalnych problemów związanych z błędami czujników czy funkcjonalnością oprogramowania.

Komponenty wewnętrzne składają się z pompy i elektrozaworu, które są sterowane przez mikrokontroler z informacjami zwrotnymi od czujnika ciśnienia powietrza. W trakcie pracy elektrozawór jest zamknięty, podczas gdy pompa nadmuchiwa kamizelkę lub inny element odzieży do pożądanego ciśnienia. Po upływie zaplanowanego czasu następuje deflacja odzieży i cykl się powtarza. Podczas gdy cykliczny efekt ciśnienia można uzyskać za pomocą ręcznego zaworu w kamizelce, zautomatyzowanie urządzenia sprawia, że system jest znacznie bardziej praktyczny przy dłuższych okresach – używanie go w ten sposób byłoby trudne z ręczną pompką. Urządzenie to może również pomóc w nadmuchiwaniu dla osób, które nie mają siły lub zręczności, aby używać ręcznej pompy.

Autor planuje, oprócz prostego interfejsu haptycznego, wzbogacić system o algorytm uczenia maszynowego dla potencjalnej aktywacji za pomocą różnych danych z innych sensorów. Może to zapewnić dodatkową metodę aktywacji dla osób z autyzmem lub będzie

to sposób na uruchamianie systemu w trakcie szczególnie kryzysowych momentów. Jak to będzie działać i jakie rodzaje czujników zostaną użyte, wciąż jest jeszcze ustalane.

Autor zastosował w systemie zasilanie z akumulatora oraz prostej przetwornicy 3,3 V. System kontrolowany jest przez mikrokontroler ESP32-S3. Układ ten będzie wystarczająco wydajny zarówno do działań wewnętrznych, jak i potencjalnej dalszej integracji z zewnętrznym pakietem czujników.

Aktualnie przetestowana pompka i rozwiązanie baterijne mogą nadmuchać kamizelkę do odpowiedniego nacisku w około 10...15 sekund. To wywołuje efekt immobilizacji i powinno być stosowane tylko przez krótki czas. Niższe ciśnienie można osiągnąć w mniej niż 10 sekund, co pozwala na dłuższe noszenie kamizelki.

Gdy sprzęt i oprogramowanie dla podstawowej funkcjonalności zostaną ukończone, wszystkie pliki projektowe mają zostać udostępnione jako otwarte dla innych użytkowników, w tym możliwe wersje płytki PCB do różnych zastosowań.

<https://hackaday.io/project/191076-auto-inflate-ai>



Monitoring stanu drzwi do sypialni

Celem tego projektu było opracowanie systemu, który mógłby poinformować personel domu opieki o przypadkach pacjentów z demencją wchodzących do pokoi innych osób, zwłaszcza gdy druga osoba znajduje się w nocy w łóżku.

System składa się z jednostki bazowej na biurku pielęgniarek, która zawiera zegar czasu rzeczywistego, rejestrator danych na kartę SD, wyświetlacz i brzęczyk, a także przyciski i system menu do konfiguracji. System ma również kilka przenośnych odbiorników z wyświetlaczem i brzęczykiem, które personel medyczny może nosić przy sobie, gdy nie jest w pobliżu jednostki bazowej na biurku.

System był projektowany w czasach, gdy moduły Wi-Fi kosztowały ponad 30 dolarów za sztukę! Taki system musiałby mieć ich ponad 40 sztuk, co nie było wykonalne w założonym budżecie. Został więc zastosowany moduł NRF24. Autor wskazuje, że gdyby miał go budować ponownie, użyłby modułów Wi-Fi.

Pierwotny plan zakładał użycie magnetycznych czujników otwarcia drzwi i zasilanych baterijnie modułów Arduino z NRF24 na każdych drzwiach, aby monitorować otwieranie i zamykanie drzwi. Okazało się, że istniał już czujnik ruchu w placówce, który obserwował każde wejście do pokoi, ale logika istniejącego systemu czujników łóżkowych uniemożliwiała wywołanie alarmu, jeśli mieszkaniowiec nadal był w łóżku (np. jeśli inny mieszkaniowiec wszedł do pokoju). Zastosowanie czujnika PIR zamontowanego na suficie miało kilka kluczowych zalet:

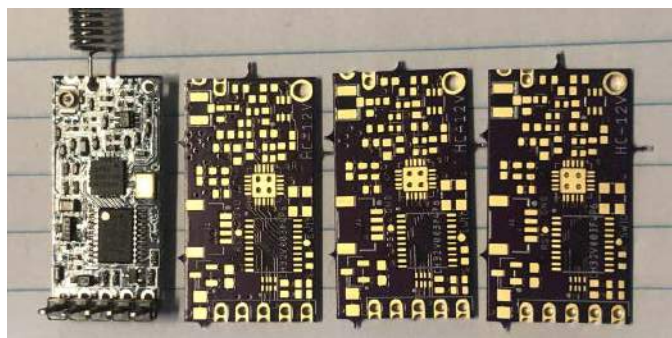
- wykrywał aktywność przy drzwiach nawet, gdy były otwarte,
- nie wymagał baterii, które mogłyby się wyczerpać,
- mniej jednostek radiowych w sieci oznaczało prostszą topologię,
- mniej komponentów sprzętowych,
- nie trzeba było instalować niczego na lub w drzwiach.

Dlatego autor uzyskał zgodę od operatora systemu czujników łóżkowych i zarządcy obiektu, aby odłączyć wszystkie istniejące czujniki ruchu od systemu czujników łóżkowych i podłączyć je do jednego z 5

hubów Arduino/NRF24, które przekazywały zdarzenia ruchu do jednostki bazowej na biurku pielęgniarzek.

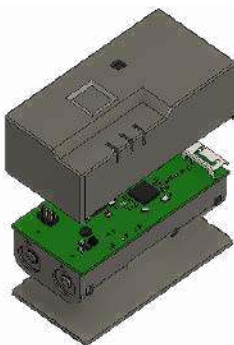
W projekcie można znaleźć dokładny opis, krok po kroku, jak system ten został zaprojektowany i zintegrowany z pozostałymi, starszymi komponentami.

<https://hackaday.io/project/191066-dsu-bedroom-door-monitor>



Mostek szeregowy HC12 433 MHz dla RISC-V

Oryginalny moduł HC-12 zawierał niedrogi 8-bitowy STM8S003F3 w połączeniu z układem firmy Silicon Labs – Si4463 i tworzył bezprzewodowy most UART, łączący dwa zdalne urządzenia na duże odległości – do 2 km. Firma WCH wprowadziła na rynek moduł CH32V003 w 2023 roku. Układ ten ma prawie identyczny układ pinów jak mikrokontroler STM8, ale zawiera 32-bitowy rdzeń RISC-V, dwukrotnie większą pamięć Flash i pamięć RAM. Ten projekt ma na celu ulepszenie oryginalnego modułu HC-12 i dodanie kilku przydatnych opcji poprzez zmianę mikrokontrolera sterującego.



<https://tiny.pl/cpvk2>



Prosta gra Cyclone z pierścieniem diod WS2812

Ta gra bazuje na automacie do gier Cyclone, w której gracz próbuje zatrzymać diodę LED poruszającą się wokół koła w określonym miejscu. Dodatkowo gra ma wyświetlacz LCD, który pokazuje najwyższy wynik i bieżącą rundę.

Kod programu bazuje na oprogramowaniu napisanym przez Joerna Weise. Autor pobrał je z jego repozytorium na GitHubie i zmodyfikował. Skrócił on sposób testowania LED-ów, a następnie dodał różne dźwięki dla każdego segmentu gry, dzięki czemu jest teraz znacznie bardziej interesująca.

Urządzenie jest bardzo proste w budowie i składa się z modułu Arduino Nano, pierścienia z 12 diodami LED WS2812,

wyświetlacza LCD 16x2 z protokołem komunikacji I²C, dwóch przycisków i brzęczyka. Co do samej rozgrywki, jak już wspomniano, w tym przypadku nie ma poziomów ze zwiększającą się prędkością, ale każda kolejna runda rozpoczyna się z losową prędkością, a ogólnie prędkości można łatwo zmienić w kodzie.

Po włączeniu gry wszystkie diody LED zapalają się sekwencyjnie, akompaniowane odpowiednim efektem dźwiękowym, a na wyświetlaczu LCD pojawia się odpowiednie powiadomienie o teście. Następnie, po naciśnięciu przycisku, gra się rozpoczyna. Celem jest naciśnięcie przycisku w momencie, gdy obracająca się dioda znajduje się dokładnie na określonych polach. W dwóch pierwszych poziomach są to trzy diody, a w kolejnych poziomach tylko jedna. Liczba przebytych okrążeń i wynik są pokazywane na wyświetlaczu. Jeśli nie uda nam się trafić w czerwoną diodę LED, gra się kończy, a na wyświetlaczu pojawia się wynik. Liczba przebytych okrążeń i wynik są pokazywane na wyświetlaczu. Najwyższy wynik jest zapisywany w pamięci EEPROM mikrokontrolera, dzięki czemu jest zachowany nawet po zresetowaniu urządzenia (jeśli chcemy usunąć wynik najwyższy, podczas włączania urządzenia należy przytrzymać przycisk HSR).

<https://tiny.pl/cpvk2>



Rękawiczka z sensorami do nadawania wiadomości alfabetem Morse'a

Projekt ten powstał na prośbę krótkofalowca, który mieszkał niedaleko autora. Na lokalnej liście e-mailowej zadał on pytanie, w jaki sposób mógłby kontynuować nadawanie CW (fala ciągła, znana również jako alfabet Morse'a) z jego problemami z poruszaniem się.

Krótkofalowiec zapytał, czy są jakieś specjalne rękawiczki do sterowania radiem. Jakkolwiek na rynku nie ma takich pomocy, to autor tego projektu nie mógł pozbyć się tej idei z głowy, dlatego też zakupił szpulkę nici przewodzącej, rękawiczkę i zabrał się do pracy. W artykule źródłowym opisuje on swoje zmagania.

<https://hackaday.io/project/191017-cw-glove>

**Podstawowe parametry:**

- dwa kanały USB-C/UART lub MPSSE,
- każdy z kanałów może zostać skonfigurowany niezależnie,
- może pełnić funkcję programatora JTAG,
- konfiguracja za pomocą oprogramowania FT_Prog i zapis ustawień w pamięci nieulotnej.

*** Uwaga!** Elektroniczne zestawy do samodzielnego montażu. Wymagana umiejętność lutowania! Podstawową wersją zestawu jest wersja [B] nazywana potocznie KIT-em (z ang. zestaw). Zestaw w wersji [B] zawiera elementy elektroniczne (w tym [UK] – jeśli występuje w projekcie), które należy samodzielnie wstawić w dotychczasową płytkę drukowaną (PCB). Wykaz elementów znajduje się w dokumentacji, która jest podlinkowana w opisie kitu. Mając na uwadze różne potrzeby naszych klientów, oferujemy dodatkowe wersje:

Dodatkowe materiały do pobrania ze strony www.ulubionykiosk.pl/media

- Moduł czterech wyjść przełącznikowych z interfejsem USB-C (EP 10/2023)
- Moduł wejść cyfrowych z optoizolacją i interfejsem USB-C (EP 9/2023)
- AVT5988 Konwerter USB-C-RS485 (EP 7/2023)
- AVT5717 Konwerter USB-UART w standardzie Grove (EP 5/2023)
- AVT5648 Konwerter USB-UART z eksterndem (EP 10/2019)
- AVT5648 Izolowana przejściówka USB/UART (EP 9/2018)

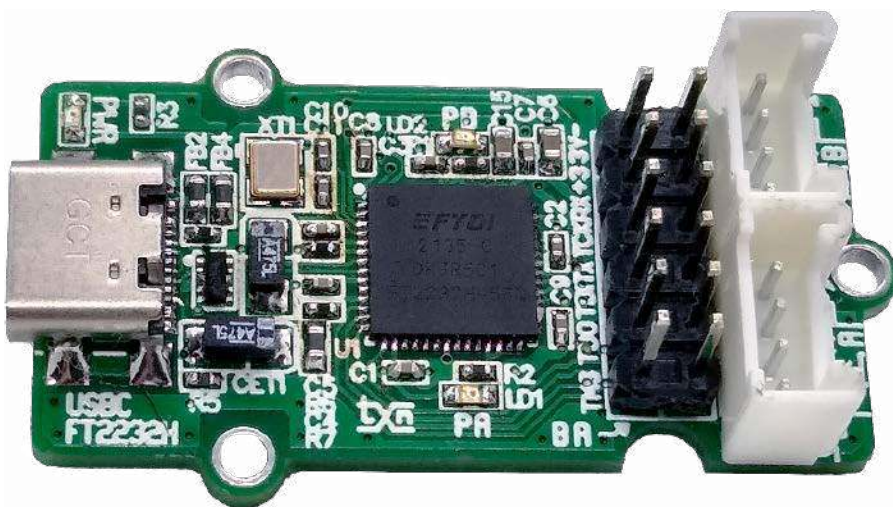
- wersja [C] – zmontowany, uruchomiony i przetestowany zestaw [B] (elementy wstawiane w płytkę PCB),
 - wersja [A] – płytkę drukowaną bez elementów i dokumentacji.
- Kity, w których występuje układ scalony wymagający zaprogramowania, mają następujące dodatkowe wersje:
- wersja [A+] – płytkę drukowaną [A] + zaprogramowany układ [UK] i dokumentacja,
 - wersja [UK] – zaprogramowany układ.

Nie każdy zestaw AVT występuje we wszystkich wersjach! Każda wersja ma załączony ten sam plik PDF! Podczas składania zamówienia upewnij się, którą wersję zamawiasz! <http://sklep.avt.pl>

W przypadku braku dostępności na stronie sklepu osoby zainteresowane zakupem płytek drukowanych (PCB) prosimy o kontakt via e-mail: kity@avt.pl.

Dwukanałowy konwerter USB-C z układem FT2232H

Dwukanałowy konwerter USB-C, który zawiera układ FT2232H, może pełnić funkcję podwójnego portu UART lub w trybie MPSSE emulować szybkie interfejsy JTAG, SPI, I²C. Druga cecha jest szczególnie przydatna, gdyż FT2232H może wtedy pełnić funkcję programatora JTAG. Umożliwia to budowę niskobudżetowego programatora przydatnego np. do kursu FPGA Lattice publikowanego na łamach EP.



Układ FT2232H pełni funkcje dwukanałowego konwertera USB, a realizowana funkcja konfigurowana jest oprogramowaniem FT_Prog i zapisywana w pamięci nieulotnej typu 93LC66B. Każdy z kanałów może zostać skonfigurowany niezależnie, co umożliwia realizację kanałów JTAG i UART przydatnych zarówno do programowania, jak i testowania aplikacji. FT2232H pracuje w typowej aplikacji zasilanej z portu USB.

Budowa i działanie

Konwerter zawiera układ FT2232H, którego budowę wewnętrzną pokazuje **rysunek 1**. Zawiera on wszystkie niezbędne do budowy konwertera bloki funkcjonalne. Schemat konwertera został pokazany na **rysunku 2**. Magistrala USB doprowadzona jest do złącza USB-C, pracującego w trybie zgodności z USB2.0. Matryca diod TVS zabezpiecza szynę danych DN/DP i zasilanie VBUS interfejsu USB-C przed skutkami przepięć. Stabilizator LDO U3 typu

AP7361 dostarcza napięcia 3,3 V do zasilania układu konwertera i umożliwia zasilanie układów współpracujących, podłączonych do złączy PA/B, UARTA/B.

Każdy kanał ma diodę Led (PA, PB) sygnalizującą status komunikacji. Interfejs szeregowy UART wraz z zasilaniem 3,3 V doprowadzony jest do złączy UARTA/B zgodnych z Grove, z których można korzystać, gdy FT2232 realizuje funkcje UART. W trybie MPSSE (wieloprotokołowy szybki interfejs szeregowy) konieczne jest wyprowadzenie 4 sygnałów dla każdego kanału, które wraz z zasilaniem dostępne są na złączach PA/B.

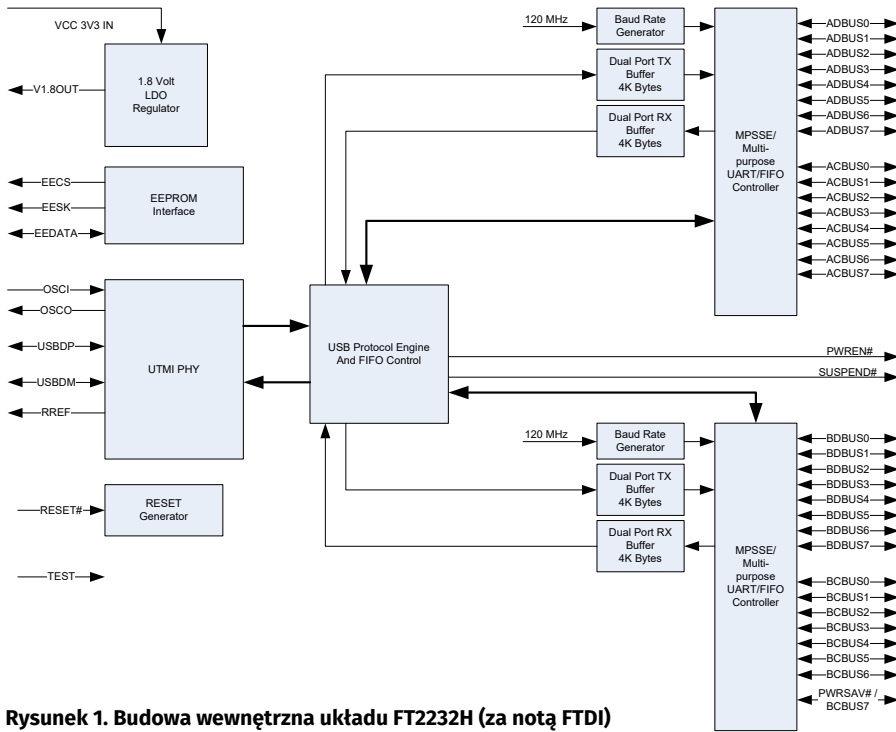
Przypisanie sygnałów zależy od wybranego interfejsu, np. dla JTAG są to TCK, TDI, TDO, TMS, dla SPI SCK, DO, DI, CS. W przypadku realizacji interfejsu JTAG należy zwrócić uwagę na noty katalogowe współpracujących układów i zapewnić odpowiednie rezystory podciągające lub pojemności obciążające poszczególne sygnały, które nie są przewidziane na płytce konwertera.

Montaż i uruchomienie

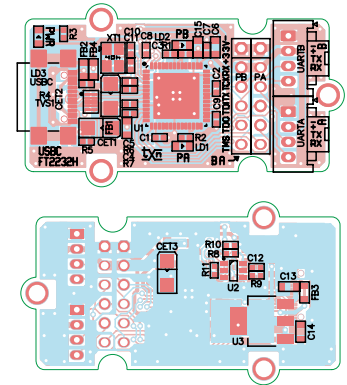
Układ zmontowany jest na miniaturowej dwustronnej płytce drukowanej, ze względu na chęć zachowania niewielkiego rozmiaru zastosowano część elementów w rozmiarze SMD0402, co wymaga sporej precyzji podczas montażu. Schemat płytki PCB został pokazany na **rysunku 3**. Montaż nie wymaga szczegółowego opisu, gotowe urządzenie zostało pokazane na fotografii tytułowej.

Po podłączeniu do komputera z systemem Windows układ wykrywany jest automatycznie. W fabrycznej konfiguracji FT2232 bez dodatkowej konfiguracji obsługuje podwójny interfejs UART. Dla sprawdzenia transmisji można połączyć ze sobą dwa kanały konwertera i przeprowadzić transmisję znakową z dwóch terminali portu szeregowego lub podłączyć układ do urządzenia docelowego i sprawdzić poprawność komunikacji oraz sygnalizację transmisji RX/TX diodami PA/PB.

Jeżeli chcemy użyć interfejsu MPSSE, musimy uruchomić konfigurator FT_prog.exe

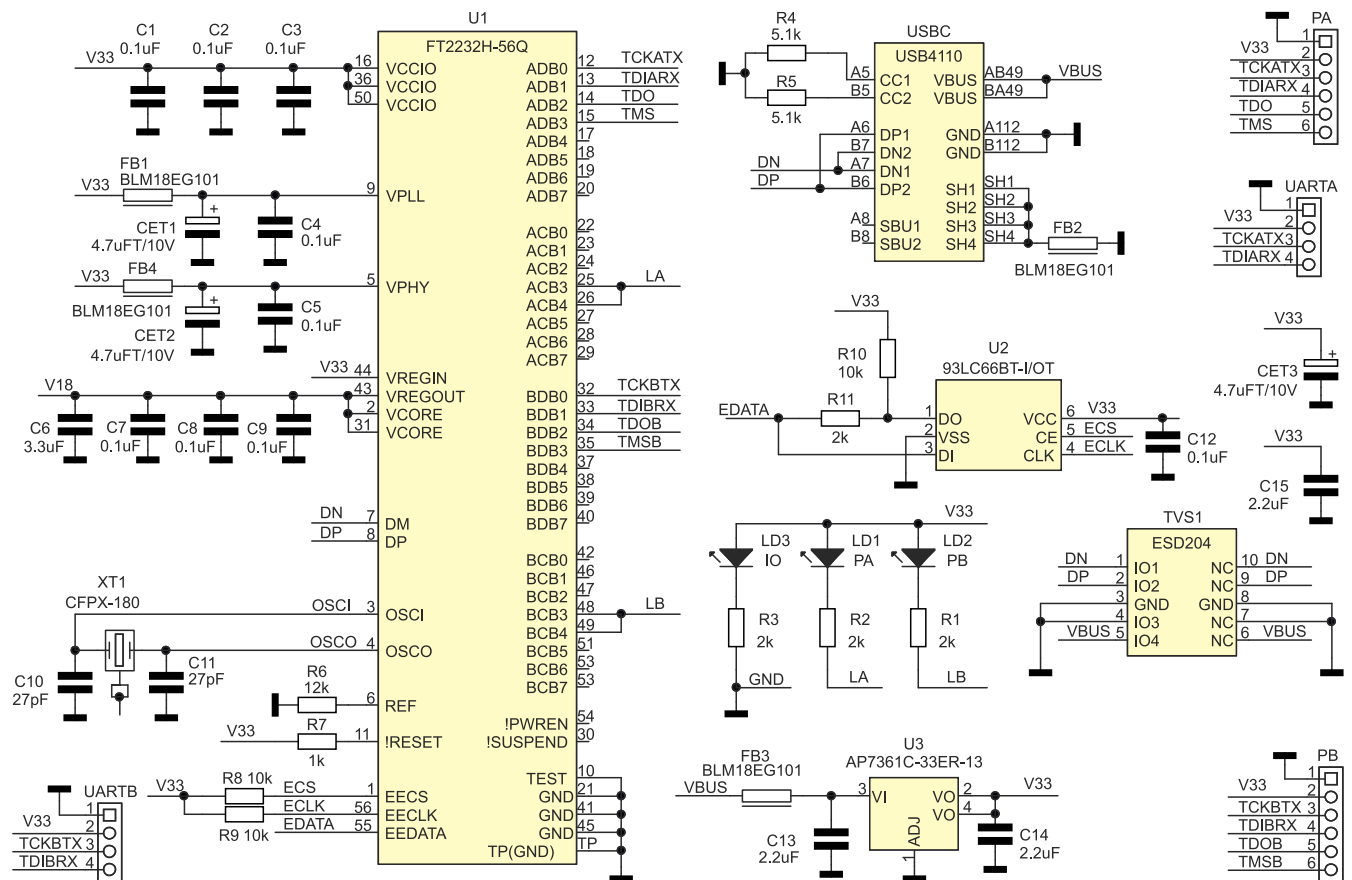


Rysunek 1. Budowa wewnętrzna układu FT2232H (za notą FTDI)



Rysunek 3. Schemat płytki PCB

(rysunek 4) i zmienić konfigurację (PortX\Hardware) z RS232UART na 245FIFO oraz aktywować driver D2XX (PortX\Driver). Warto ustawić maksymalny pobierany prąd z USB na 500 mA i zwiększyć wydajność GPIO do 16 mA. Po zapisaniu parametrów do EEPROM (ProgramDevice) i restarcie układu (CyclePort) układ zainicjuje się z nowymi ustawieniami.



Rysunek 2. Schemat konwertera

Wykaz elementów: (kupuj na stronie sklep.avt.pl lub osobiście: Warszawa, ul. Leszczyńska 11, tel. +48222578451, e-mail: handlowy@avt.pl)

Rezystory:

(SMD0402, 1%)
R1, R2, R3, R11: 2 kΩ
R4, R5: 5,1 kΩ
R6: 12 kΩ
R7: 1 kΩ
R8, R9, R10: 10 kΩ

Kondensatory:

C1, C2, C3, C4, C5, C7, C8, C9, C12: 0,1 μF 10 V (SMD0402)
CET1, CET2, CET3: 4,7 μF/10 V, tantalowy A (3216)

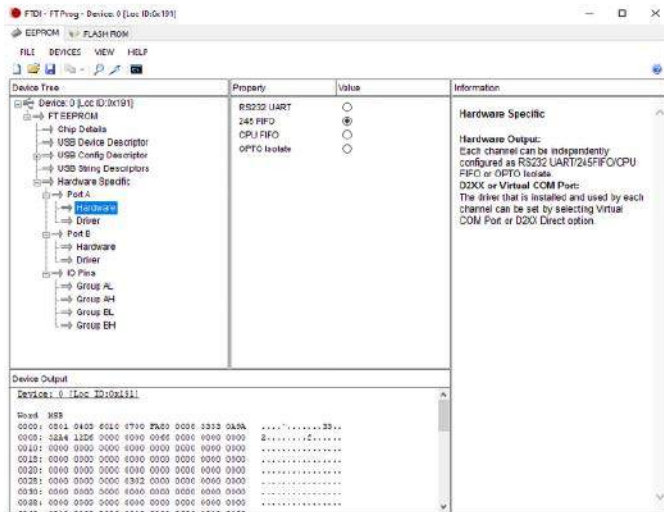
C6: 3,3 μF 10 V (SMD0603)
C10, C11: 27 pF 10 V (SMD0402)
C13, C14, C15: 2,2 μF 10 V (SMD0603)

Półprzewodniki:

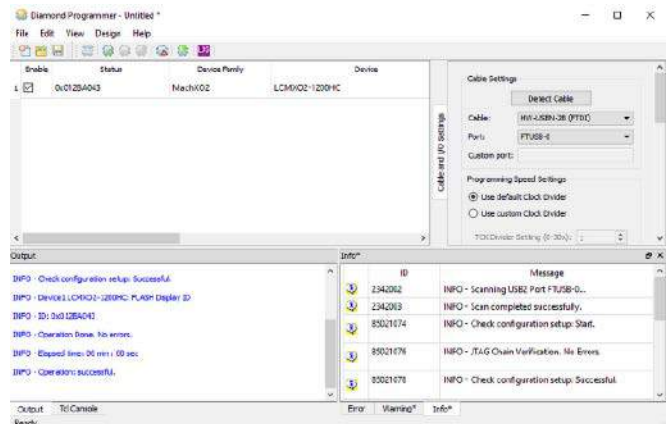
LD1, LD2: dioda LED żółta (SMD0603)
LD3: dioda LED czerwona (SMD0603)
U1: FT2232H-56Q (VQFN56)
U2: 93LC66BT-I/OT (SOT-23-6)
U3: AP7361C-33ER-13 (SOT-223)

Pozostałe:

FB1, FB2, FB3, FB4: dławik ferrytowy BLM18EG101 (SMD0603)
PA, PB: złącze szpilkowe SIP6, 2,54 mm
TVS1: ESD204 (USON10)
UARTA, UARTB: złącze Grove proste (110990030)
USB-C: złącze USB-C USB4110GTC
XT1: kwarc 12 MHz (CFPX-180)



Rysunek 4. Konfiguracja FT2322



Rysunek 5. Interfejs JTAG

W przypadku zmiany funkcji z UART układ nie jest już widoczny w zakładce portów szeregowych USB Menedżera urządzeń. Aktywny jest driver D2xx, a obsługa MPSSE musi zająć się współpracującą aplikacją, np. programator JTAG. Aby używać płytki konwertera jako programatora

do kursu FPGA Lattice, kanał A ustawiamy jako JTAG, kanał B jako interfejs szeregowy. Płytkę jest rozpoznawana jako programator zgodny z HW-USB-2B (FTDI) (rysunek 5). Należy jednak pamiętać o wyborze kanału FTUSB-0. Po podłączeniu portu A konwertera z interfejsem JTAG układu i inicjacji

łańcucha JTAG można wykryć podłączony układ, w tym przypadku LCMXO2-1200HC i używać programatora jak „fabrycznego”. Kanału UART można użyć do komunikacji z aplikacją docelową.

Adam Tatuś, EP

REKLAMA

sklep.avt.pl

- Nauka elektroniki
- AVT Kits
- Elektronika
- Sprzęt pomiarowy i zasilanie
- Warsztat
- Dom i ogród





W ofercie AVT*

AVT6009

Podstawowe parametry:

- liczba obsługiwanych kart RFID/NFC: 100,
- liczba obsługiwanych PIN-kodów: 100,
- zasilanie: 8...10 VDC, 150 mA,
- parametry styków przełącznika: 250 VAC, 125 VDC, AC: 10 A; DC: 8 A (szczegóły w dokumentacji przełącznika).

* **Uwaga!** Elektroniczne zestawy do samodzielnego montażu. Wymagana umiejętność lutowania! Podstawową wersją zestawu jest wersja [B] nazywana potocznie KIT-em (z ang. zestaw). Zestaw w wersji [B] zawiera elementy elektroniczne (w tym [UK] – jeśli występuje w projekcie), które należy samodzielnie wlutować w dołączoną płytkę drukowaną (PCB). Wykaz elementów znajduje się w dokumentacji, która jest podlinkowana w opisie kitu. Mając na uwadze różne potrzeby naszych klientów, oferujemy dodatkowe wersje:

- **wersja [C]** – zmontowany, uruchomiony i przetestowany zestaw [B] (elementy wlutowane w płytkę PCB),
 - **wersja [A]** – płytką drukowaną bez elementów i dokumentacji.
- Kity, w których występuje układ scalony wymagający zaprogramowania, mają następujące dodatkowe wersje:
- **wersja [A+]** – płytką drukowaną [A] + zaprogramowany układ [UK] i dokumentacja,
 - **wersja [UK]** – zaprogramowany układ.

Dodatkowe materiały do pobrania ze strony www.ulubionykiosk.pl/media

- Simple Access System 2 (2) (EP 6/2022)
- Simple Access System 2 (1) (EP 5/2022)
- NFC Lock (EP 4/2022)
- AVT5186 Bezstykowy zamek RFID
- AVT969 Bezstykowy zamek RFID
- AVT886 System bezstykowej kontroli dostępu (EP 10/2000)

Nie każdy zestaw AVT występuje we wszystkich wersjach! Każda wersja ma załączony ten sam plik PDF! Podczas składania zamówienia upewnij się, którą wersję zamawiasz!
<http://sklep.avt.pl>

W przypadku braku dostępności na stronie sklepu osoby zainteresowane zakupem płytek drukowanych (PCB) prosimy o kontakt via e-mail: kity@avt.pl.

multiLock

multiLock to system kontroli dostępu, który umożliwia dostęp poprzez karty zbliżeniowe standardu RFID/NFC oraz zwykłe 4-cyfrowe PIN-kody. W celu obsługi wspomnianych kart zastosowano moduł czytnika RFID/NFC firmy NXP pod postacią peryferium oznaczonego symbolem PN532, który integruje w sobie specjalizowany układ o tej samej nazwie (mikrokontroler z rdzeniem 80C51 wyposażony w 40 kB pamięci ROM i 1 kB pamięci RAM) oraz cały interfejs RF.

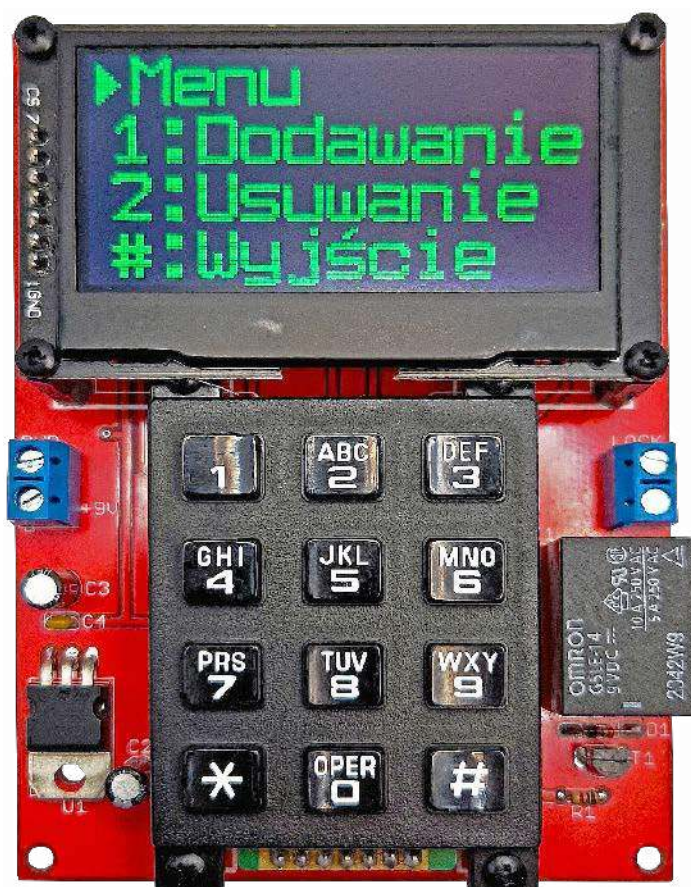
W portfolio moich projektów bez problemu znajdziecie różnego rodzaju systemy dostępu, w których medium autoryzacyjnym były pastylki iButton, karty RFID/NFC czy zwykłe PIN-kody, jednak przyznać muszę, że zawsze były to systemy, w implementacji których zdecydowałem się na zastosowanie wyłącznie jednej ze wspomnianych technologii. Nie powinno to specjalnie dziwić, gdyż dokładnie w ten sposób implementowane są zazwyczaj komercyjne rozwiązania tego typu, jednak musicie przyznać, że stanowi to swego rodzaju ograniczenie.

Podam prosty przykład. Załóżmy, że korzystamy z systemu kontroli dostępu, który stosuje karty RFID/NFC w celu autoryzacji dostępu. Czy w przypadku, gdy nie mamy przy sobie tego rodzaju karty, istnieje możliwość autoryzowanego dostępu? Oczywiście, że NIE! Czy tego rodzaju rozwiązanie jest satysfakcjonujące? Oczywiście, że NIE! Jak powinno się rozwiązać tego rodzaju zagadnienie? Moim zdaniem, każdy tego rodzaju system powinien zapewniać kilka mechanizmów kontrolowanego dostępu zawierających różne media autoryzacyjne. Nie ukrywam, że do wniosków takich doszedłem, używając jeden z nowszych, komercyjnych systemów kontroli dostępu, gdzie medium autoryzacyjnym były zarówno karty/breloki standardu RFID/NFC, jak i zwykłe PIN-kody...

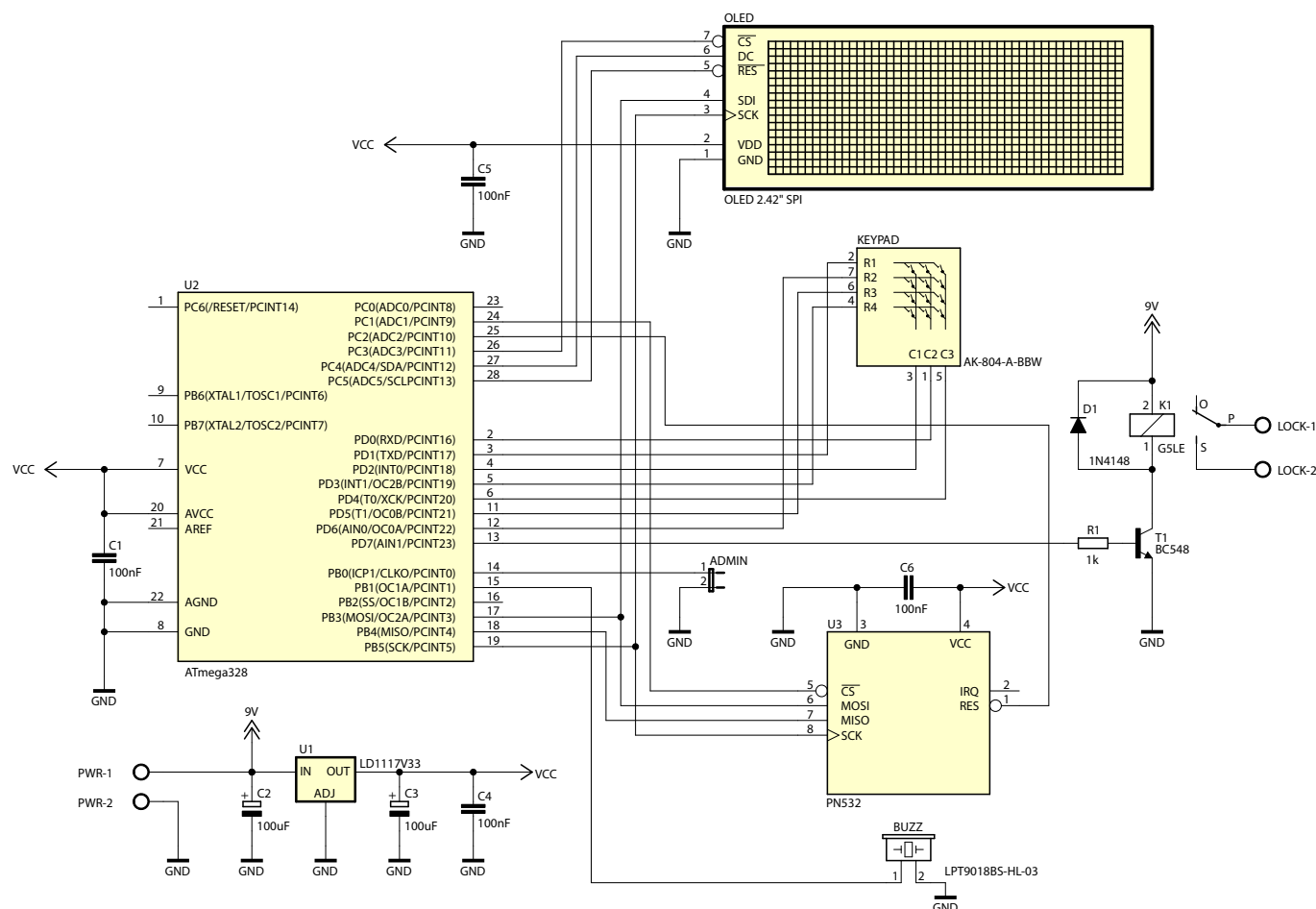
i właśnie tego rodzaju implementację chciałem Wam zaprezentować w ramach niniejszego projektu.

Moduł PN532

Moduł, o którym już wspomniałem, to podzespół pozwalający na wymianę danych (jednokierunkową lub dwukierunkową) z kartami i urządzeniami w standardzie RFID i NFC. Obsługiwane rodzaje kart to: Mifare 1k, 4k, Ultralight, DesFire, ISO/IEC 14443-4, Innovision Jewel (IRT5001) oraz FeliCa (RCS 860, RCS 854). Oprócz obsługi kart lub tagów RFID, moduł pozwala również na dwukierunkową komunikację z urządzeniami NFC – np. telefonem komórkowym, a także wymianę danych między dwoma modułami w trybie P2P.



Wspomniany podzespół wyposażono dodatkowo w wygodne interfejsy komunikacyjne, dzięki którym możliwa jest jego konfiguracja i wymiana danych z procesorem hosta. Interfejsy te to: I²C, SPI lub HSU (*High Speed UART*). Wybór aktywnego interfejsu komunikacyjnego dokonywany jest za pomocą dedykowanych switchy SMD zintegrowanych na obwodzie drukowanym PCB. Aby przygotować moduł do pracy z interfejsem SPI, przełączamy przełącznik oznaczony jako ON w pozycję On (blisko znacznika „1”), zaś przełącznik oznaczony jako KE w pozycję Off (blisko znacznika „KE”). Na tym samym obwodzie drukowanym zintegrowano również antenę poprowadzoną na brzegach płytki, dzięki czemu dostajemy urządzenie gotowe do pracy. Tyle w kwestii samego modułu.



Rysunek 1. Schemat ideowy urządzenia multiLock

Kilka słów uwagi należy w tym miejscu poświęcić samej metodologii przesyłania danych RFID/NFC. Obie to dwie bardzo zbliżone do siebie technologie bezprzewodowej komunikacji, które używane są praktycznie na każdym kroku w różnych dziedzinach naszego życia – od weryfikacji tożsamości poprzez kontrolę dostępu aż do systemów bezprzewodowych płatności zbliżeniowych. Jakże są główne różnice pomiędzy RFID a NFC? RFID pozwala na jednokierunkową (w sensie inicjowania transferu), bezprzewodową komunikację pomiędzy niezasilonym znacznikiem a zasilanym czytnikiem RFID, przy czym odległość między tym czytnikiem a znacznikiem może dochodzić nawet do 200 m (w wybranych implementacjach) i zależy w dużej mierze od zastosowanych zakresów częstotliwości radiowych oraz protokołów komunikacji.

W odróżnieniu od RFID, technologia NFC pozwala natomiast na dwustronną komunikację pomiędzy dwoma urządzeniami NFC, realizując obsługę bardziej złożonych

mechanizmów, a także wymianę danych. Wynika to głównie z idei powstania tej technologii, która w zamierzeniu miała umożliwiać komunikację na małe, bezpieczne odległości (poniżej 20 cm), przez co idealnie nadaje się do zastosowania w telefonach komórkowych do przeprowadzania bezpiecznych transakcji bezgotówkowych lub wymiany danych pomiędzy tymi urządzeniami.

W prezentowanym urządzeniu skupię się wyłącznie na obsłudze kart Mifare 1k oraz urządzeń NFC zgodnych z tym standardem, gdyż naszym zadaniem będzie wyłącznie odczytanie unikalnego numeru seryjnego karty/urządzenia RFID/NFC (tzw. UID). Należy mieć jednak świadomość, że karty tego rodzaju udostępniają dużo większą funkcjonalność. Dla przykładu wyposażono je w 1024 bajty pamięci EEPROM (16 sektorów po 4 bloki o rozmiarze 16 B), która może być zabezpieczona przed zapisem/odczytem za pomocą zdefiniowanego klucza dostępowego (i to indywidualnie dla każdego z tychże sektorów).

Ciekawe, nieprawdaż? I na tym zakończę opis samego modułu i zagadnień związanych z obsługą kart RFID/NFC, mimo że jest to naprawdę dość obszerne i niezmiernie ciekawe zagadnienie. Dlaczego? Otóż dlatego, że tematyka ta została przeze mnie szczegółowo opisana w ramach projektu o nazwie „NFC Lock”, który opublikowano na łamach naszego miesięcznika w wydaniu EP 04/22 (KIT AVT 5926) [1], w związku z czym zainteresowanych Czytelników odsyłam do tych materiałów. Jeśli zaś chodzi o nasz niniejszy projekt, to, jak napisałem na wstępie, w celu autoryzacji dostępu używa on alternatywnego sposobu w postaci zwykłego, 4-cyfrowego kodu PIN, którego wprowadzenie możliwe jest dzięki zastosowaniu standardowej klawiatury matrycowej. Tyle tytułem wstępu.

Budowa i działanie

Przejdźmy do schematu ideowego naszego urządzenia, który pokazano na **rysunku 1**. Jak widać, zaprojektowano bardzo prosty system mikroprocesorowy,

Wykaz elementów: (kupuj na stronie sklep.avt.pl lub osobiście: Warszawa, ul. Leszczyńska 11, tel. +48222578451, e-mail: handlowy@avt.pl)

Półprzewodniki:

U1: LD1117V33 (TO-220)
U2: ATmega328 (DIP28)
U3: moduł RFID/NFC typu PN532
OLED: wyświetlacz OLED 2.42", kontroler SSD1309, interfejs SPI
D1: 1N4148 (DO-35)
T1: BC548 (TO-92)

Rezystory:

R1: 1 kΩ (miniaturowy 1/8 W, raster 0,2")

Kondensatory:

C1, C4, C5, C6: 100 nF ceramiczny (raster 0,1")
C2, C3: 100 μF/16 V elektrolityczny (raster 0,1")

Pozostałe:

K1: przekaźnik G5LE-14-9
PWR, LOCK: złącze śrubowe AK500/2 (raster 0,1")
KEYPAD: klawiatura numeryczna typu AK-804-A-BBW
BUZZ: przetwornik piezoelektryczny SMD typu LPT9018BS-HL-03-4.0-12-R
ADMIN: jumper THT

którego sercem jest bardzo popularny mikrokontroler firmy Microchip (dawniej Atmel) o oznaczeniu ATmega328 taktowany wewnętrznym, wysokostabilnym generatorem RC o częstotliwości 1 MHz. Mikrokontroler ten odpowiedzialny jest za sprzętową obsługę interfejsu SPI, dzięki któremu realizuje współpracę zarówno ze wspomnianym wcześniej modułem RFID/NFC, jak i bardzo efektywnym wyświetlaczem OLED o przekątnej 2,42" i rozdzielczości 128×64 piksele wyposażonym w sterownik ekranu typu SSD1309 (zbliżony do SSD1306) stanowiącym element graficznego interfejsu użytkownika.

I już w tym miejscu należy się kilka słów wyjaśnienia. Jak widzicie, oba peryferia (moduł RFID/NFC i wyświetlacz OLED) korzystają z tego samego, szeregowego interfejsu

SPI, zaś wybór aktywnego układu, dla którego przeznaczone są transmitowane dane (lub z którego odczytywane są dane – w przypadku modułu RFID/NFC), odbywa się tradycyjnie dzięki odrębnym wyprowadzeniom CS obu układów. Jest jednak pewien istotny haczyk. Otóż moduł RFID/NFC komunikuje się z mikrokontrolerem dzięki transmisji szeregowej SPI w kolejności od bitu najmłodszego (LSB), zaś wyświetlacz OLED w kolejności od bitu najstarszego (MSB). W takim wypadku każdorazowa zmiana docelowego peryferium wymaga reinicjalizacji modułu SPI po stronie mikrokontrolera. To bardzo istotny „drobiazg”, który łatwo pominąć, jeśli nie studujemy dokładnie dokumentacji producenta, co z kolei uniemożliwiłoby nam nawiązanie poprawnej komunikacji ze wspomnianymi wcześniej podzespołami.

Przejdźmy zatem do naszego modułu OLED. Zastosowanie wyświetlacza tego typu wynikało z chęci wyposażenia naszego urządzenia w efektywny, a zarazem bardzo czytelny interfejs użytkownika, co nie jest bez znaczenia w przypadku obsługi urządzeń tego typu. Co istotne, wyświetlacz, o którym mowa, ma bardzo przystępną cenę na dalekowschodnim portalu aukcyjnym, która oscyluje w granicach 10 \$, a dodatkowo dostępny jest w kilku możliwych kolorach, przez co łatwo dopasować docelowy design urządzenia do naszych upodobań czy wymagań klienta. Jedyne, na co należy zwrócić uwagę, to położenie i rozkład wyprowadzeń, gdyż w tym zakresie mogą występować nieznaczne różnice wynikające z kolejnych wersji wspomnianego peryferium.

Zastosowanie mikrokontrolera tego rodzaju wyposażonego w dość sporą wielkość pamięci Flash (jak na skromne AVR-y) nie wynikało z objętości programu obsługi aplikacji (który de facto zajmuje jedynie około 30% tej pamięci, z czego zdecydowana większość to tablice czcionek ekranowych i stosowanych piktogramów), a z niezbędnej wielkości zintegrowanej pamięci EEPROM, w której będziemy przechowywać dane autoryzacyjne użytkowników (ich kody PIN i numery seryjne kart RFID/NFC).

Wróćmy zatem ponownie do naszego schematu ideowego. Mikrokontroler poza zadaniami wspomnianymi powyżej realizuje obsługę klawiatury matrycowej stanowiącej element interfejsu użytkownika, jak i obsługę buzzera piezoelektrycznego niewyposażonego w zintegrowany generator, a to dzięki wykorzystaniu wbudowanego w mikrokontroler układu czasowo-licznikowego Timer1 pracującego w trybie CTC i generującego przebieg o częstotliwości 4 kHz na wyprowadzeniu PB1 (OC1A) układu. Ponadto serce naszego systemu steruje pracą przekaźnika K1 (poprzez stopień tranzystorowy T1), który w zamierzeniach sterować ma elektromechanicznym rygłem pozwalającym na dostęp do chronionego obiektu, jak i obsługując zworkę ADMIN, zadaniem której jest wprowadzenie systemu multiLock w tryb wprowadzania kodu/karty RFID/NFC administratora.

Implementacja takiej sprzętowej zwory była intencjonalna a wynikała z chęci zabezpieczenia systemu przed przypadkowym wejściem w tryb wprowadzania kodu/karty RFID/NFC administratora, wszak tego rodzaju media autoryzacyjne umożliwiają w dalszej kolejności pełną kontrolę nad listą dostępnych użytkowników i ich danych autoryzacyjnych. Co więcej, aktywna zworka ADMIN umożliwia wykasowanie całej pamięci użytkowników (poza danymi autoryzacyjnymi Administratora).

Tyle w kwestii schematu ideowego, w związku z czym przejdźmy do zagadnień programistycznych. Tak jak wspomniano wcześniej, nie będę tutaj powielał informacji

Listing 1. Plik nagłówkowy modułu obsługi klawiatury matrycowej

```
//Definicje portu klawiatury
#define KEYPAD_PORT PORTD
#define KEYPAD_PIN PIND
#define KEYPAD_DDR DDRD
#define COL1 (1<<PD2)
#define COL2 (1<<PD0)
#define COL3 (1<<PD4)
#define ROW1 (1<<PD1)
#define ROW2 (1<<PD6)
#define ROW3 (1<<PD5)
#define ROW4 (1<<PD3)

//Definicja stałej nieprzyciśniętego klawisza
#define NOT_PRESSED 255

//Definicja znaków specjalnych keypad-a
#define CODE_GWIAZDKA 10
#define CODE_HASH 11
```

Listing 2. Funkcja odpowiedzialna za inicjalizację obsługi klawiatury matrycowej

```
void keypadInit(void) {
    //Podciągnięcie wierszy klawiatury pod VCC
    KEYPAD_PORT |= ROW1|ROW2|ROW3|ROW4;
}
```

Listing 3. Definicje stałych modułu obsługi klawiatury matrycowej

```
//Tablica kolumn
const uint8_t Cols[3] = {COL1, COL2, COL3};
//Tablica wierszy
const uint8_t Rows[4] = {ROW1, ROW2, ROW3, ROW4};
//Tablica zwracanych kodów
const uint8_t Codes[12] = {1, 2, 3, 4, 5, 6, 7, 8, 9, CODE_GWIAZDKA, 0, CODE_HASH};
```

Listing 4. Funkcja odpowiedzialna za odczyt kodu wciśniętego klawisza klawiatury matrycowej

```
uint8_t keypadRead(void) {
    uint8_t Key = NOT_PRESSED;

    //Ustalamy numer wciśniętego klawisza
    for(uint8_t Col=0; Col<3; Col++) {
        //Wybraną kolumnę ustawiamy, jako wyjściową ze stanem 0
        KEYPAD_DDR |= Cols[Col];
        //Szukamy czy nie zwarto jej do wybranego wiersza
        for(uint8_t Row=0; Row<4; Row++) {
            if(!(KEYPAD_PIN & Rows[Row])) {
                Key = Col + (Row*3);
                break;
            }
        }
        //Wybraną kolumnę ustawiamy, jako wejściową
        KEYPAD_DDR &= ~Cols[Col];
    }
    _delay_ms(30);

    //Wszystkie kolumny ustawiamy jako wyjścia ze stanem 0,
    //żeby sprawdzić czy nadal naciskamy jakiś klawisz
    KEYPAD_DDR |= COL1|COL2|COL3;
    //Sprawdzamy czy jakikolwiek klawisz nadal jest naciśnięty
    if((KEYPAD_PIN & ROW1) && (KEYPAD_PIN & ROW2)
        && (KEYPAD_PIN & ROW3)
        && (KEYPAD_PIN & ROW4)) Key = NOT_PRESSED;
    else if(Key != NOT_PRESSED) Key = Codes[Key];
    //Wszystkie kolumny ustawiamy jako wejścia
    KEYPAD_DDR &= ~(COL1|COL2|COL3);

    return Key;
}
```

Listing 5. Funkcja inicjalizacyjna sterownika ekranu OLED o przekątnej 2.42" wyposażonego w sterownik ekranu SSD1309

```

void OLEDinit(void) {
    //Wszystkie porty jako wyjściowe
    //CS w stanie "1" (inactive), a RES w stanie "0" (reset)
    OLED_PORT |= (1<<OLED_CS);
    OLED_DDR |= (1<<OLED_DC)|(1<<OLED_CS)|(1<<OLED_RES);

    _delay_ms(10);
    OLED_SET_RES;
    _delay_ms(100);

    SSD1309writeCmd(DISPLAY_OFF_CMD);
    SSD1309writeCmd(SET_START_LINE_CMD | 0x00);

    SSD1309writeCmd(SET_CONTRAST_CMD);
    SSD1309writeCmd(0x32); //Contrast

    SSD1309writeCmd(SET_SEG_REMAPING_CMD | 0x01);
    SSD1309writeCmd(SET_COM_SCAN_REMAPPED_CMD);
    SSD1309writeCmd(NORMAL_DISPLAY_CMD);

    SSD1309writeCmd(SET_MULTIPLEX_RATIO_CMD);
    SSD1309writeCmd(0x3F); //1/64 duty

    SSD1309writeCmd(SET_DISPLAY_OFFSET_CMD);
    SSD1309writeCmd(0x00); //Not offset

    SSD1309writeCmd(SET_DISPLAY_CLOCK_DIV_CMD);
    SSD1309writeCmd(0xA0);

    SSD1309writeCmd(SET_PRECHARGE_PERIOD_CMD);
    //Set Pre-Charge as 15 Clocks & Discharge as 1 Clock
    SSD1309writeCmd(0xF1);

    SSD1309writeCmd(SET_COM_PINS_CMD);
    SSD1309writeCmd(0x12);

    SSD1309writeCmd(SET_VCOMH_DESELECT_CMD);
    SSD1309writeCmd(0x30); //Set VCOM Deselect Level

    SSD1309writeCmd(MEMORY_MODE_CMD);
    SSD1309writeCmd(0x00); //Horizontal addressing mode

    SSD1309writeCmd(CHARGE_PUMP_CMD);
    SSD1309writeCmd(0x14); //Disable

    SSD1309writeCmd(DISPLAY_ALLON_RESUME_CMD);
    SSD1309writeCmd(NORMAL_DISPLAY_CMD);
    SSD1309writeCmd(DISPLAY_ON_CMD);

    OLEDcls();
}

```

Listing 6. Plik nagłówkowy modułu obsługi wyświetlacza OLED 2,42"

```

//OLED RESOLUTION
#define OLED_WIDTH 128
#define OLED_HEIGHT 64

//OLED port definitions
#define OLED_PORT PORTC
#define OLED_DDR DDRC

#define OLED_CS PC3 //CHIP SELECT
#define OLED_DC PC4 //DATA(1)/COMMAND(0)
#define OLED_RES PC5 //RESET

#define OLED_SET_CS OLED_PORT |= (1<<OLED_CS)
#define OLED_RESET_CS OLED_PORT &= ~(1<<OLED_CS)

//DDRAM MODE
#define OLED_SET_DC OLED_PORT |= (1<<OLED_DC)
//COMMAND MODE
#define OLED_RESET_DC OLED_PORT &= ~(1<<OLED_DC)

#define OLED_SET_RES OLED_PORT |= (1<<OLED_RES)
#define OLED_RESET_RES OLED_PORT &= ~(1<<OLED_RES)

//OLED Commands
#define SET_CONTRAST_CMD 0x81
#define DISPLAY_ALLON_RESUME_CMD 0xA4
#define DISPLAY_ALLON_CMD 0xA5
#define NORMAL_DISPLAY_CMD 0xA6
#define INVERSE_DISPLAY_CMD 0xA7
#define DISPLAY_OFF_CMD 0xAE
#define DISPLAY_ON_CMD 0xAF

#define SET_DISPLAY_OFFSET_CMD 0xD3
#define SET_COM_PINS_CMD 0xDA

#define SET_VCOMH_DESELECT_CMD 0xDB

#define SET_DISPLAY_CLOCK_DIV_CMD 0xD5
#define SET_PRECHARGE_PERIOD_CMD 0xD9

#define SET_MULTIPLEX_RATIO_CMD 0xA8

#define SET_START_LINE_CMD 0x40

#define MEMORY_MODE_CMD 0x20
#define SET_COLUMN_ADDR_CMD 0x21
#define SET_PAGE_ADDR_CMD 0x22

#define SET_COM_SCAN_NORMAL_CMD 0xC0
#define SET_COM_SCAN_REMAPPED_CMD 0xC8

#define SET_SEG_REMAPING_CMD 0xA0
#define CHARGE_PUMP_CMD 0x8D

```

na temat obsługi modułu RFID/NFC, gdyż takowe są bardzo obszerne i zebrano je w ramach mojego wcześniejszego artykułu, o którym wspomniałem powyżej, a skupię się na elementach nowych.

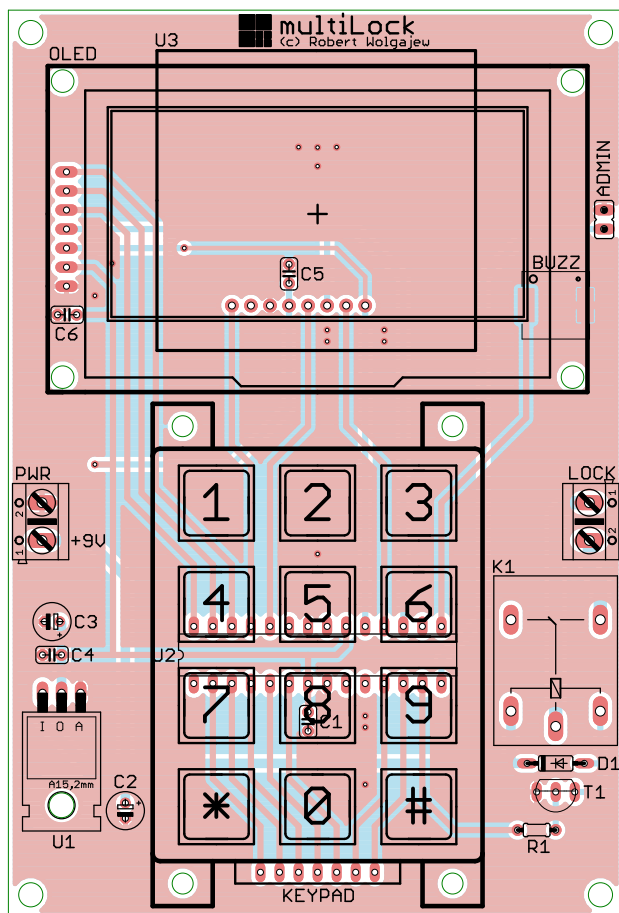
Tym nowym elementem jest obsługa klawiatury matrycowej. W sumie to nic trudnego, a już na pewno nic odkrywczego, ale pokażę Wam, jak w prosty sposób rozwiązałem tego rodzaju zagadnienie, zaznaczając, że nie jest to rozwiązanie uniwersalne, a szyte na miarę niniejszego projektu, gdzie nie ma potrzeby detekcji faktu wciśnięcia kilku klawiszy klawiatury w jednym czasie. Zaczniemy od pliku nagłówkowego do obsługi tegoż peryferium upraszczającego i porządkującego późniejszy kod, który to plik pokazano na **listingu 1**.

Dalej, na **listingu 2**, pokazano ciało funkcji odpowiedzialnej za inicjalizację obsługi klawiatury matrycowej, której jedynym zadaniem jest podciągnięcie stosownych portów wejściowych do plusa zasilania (przez rezystory podciągające zlokalizowane wewnątrz struktury mikrokontrolera).

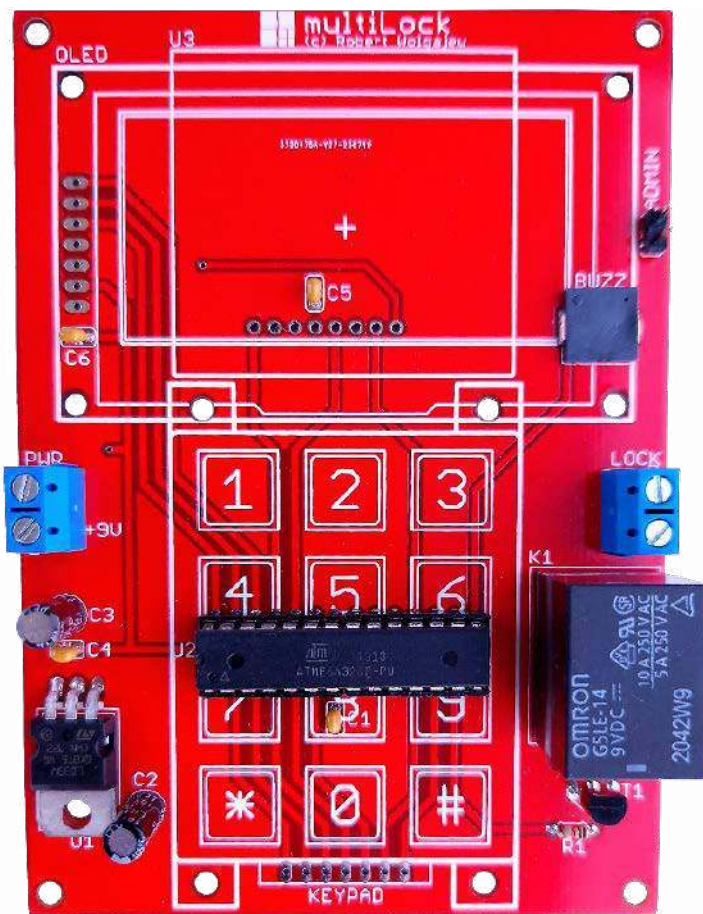
Zanim przejdę do funkcji pozwalającej na odczyt kodu przyciśniętego przycisku, niezbędne jest wprowadzenie kilku zmiennych (a w zasadzie stałych) upraszczających ten mechanizm. Stałe te zadeklarowano w pamięci RAM, co może wydawać się pewnym marnotrawstwem tej pamięci, ale nasze urządzenie korzysta z niewielkiej ilości pamięci RAM, w którą wyposażony jest mikrokontroler, przez co zadeklarowanie tych stałych w tej pamięci (zamiast w pamięci Flash) nie wydaje się jakimś kardynalnym błędem, a z pewnością przyspiesza dostęp do ich wartości. Definicje wspomnianych stałych pokazano na **listingu 3**.

I na koniec zagadnień implementacyjnych związanych z obsługą klawiatury matrycowej funkcja odpowiedzialna za odczyt kodu wciśniętego klawisza klawiatury matrycowej, której ciało pokazano na **listingu 4**.

Tyle w kwestii samej klawiatury. Zanim jednak przejdę do zagadnień montażowych, kilka niezbędnych słów uwagi na temat zastosowanego wyświetlacza OLED. Zgodnie z tym, co napisałem wcześniej, wyświetlacz nasz wyposażono w sterownik ekranu typu SSD1309 produkcji firmy Solomon Systech Limited, który w swojej konstrukcji jest bardzo zbliżony do dobrze znanego i chętnie używanego przez amatorów (jak i naszych chińskich przyjaciół) sterownika SSD1306. Pisząc bardzo zbliżony, mam na myśli, że sposób obsługi obu układów jest identyczny, ale różnią się one nieco możliwościami sprzętowymi, jak i procedurą inicjalizacji. Jako że w Internecie znajduje się mnóstwo przykładów obsługi obu sterowników, ograniczę się wyłącznie do prezentacji kodu funkcji inicjalizacyjnej, która jest charakterystyczna (w sensie wysyłanych ustawień)



Rysunek 2. Schemat montażowy urządzenia multiLock



Fotografia 1. Wygląd obwodu drukowanego urządzenia multiLock tuż przed przylutowaniem modułu RFID, wyświetlacza OLED i klawiatury matrycowej (w wersji prototypowej)

wyłącznie dla naszego modułu o przekątnej 2,42". Ciało wspomnianej funkcji pokazano na **listingu 5**. Aby zrozumieć działanie funkcji inicjalizacyjnej, niezbędna jest znajomość pliku nagłówkowego, którego treść pokazano na **listingu 6**. Tyle w kwestiach implementacyjnych.

Montaż i uruchomienie

Przejdźmy do schematu montażowego urządzenia, który został pokazany na **rysunku 2**. Jak widać, zaprojektowano bardzo zwartą konstrukcję obwodu drukowanego z wyłącznym zastosowaniem elementów THT (poza buzerem), po to, by całe urządzenie swoimi wymiarami nie przekraczało niezbędnego, minimalnego obszaru dla wykonania interfejsu użytkownika a jednocześnie było proste w konstrukcji. W tym celu w wybranych miejscach zastosowano montaż „piętrowy”, w związku z czym kolejność implementacji ma tutaj szczególne znaczenie.

Warto również podkreślić, że dla zminimalizowania zakłóceń, na płytce urządzenia poprowadzono obszerne pola masy po obu stronach obwodu drukowanego oraz zastosowano szereg przelotek pomiędzy nimi w celu zmniejszenia pojemności pasożytniczych.

Montaż urządzenia rozpoczynamy od przy-
lutowania mikrokontrolera (najlepiej w odpo-
wiedniej podstawce), stabilizatora napięcia

(U1), pozostałych elementów półprzewodnikowych, a na koniec elementów pasywnych (w tym przełącznika) i złączy PWR, LOCK oraz jumpera ADMIN. W dalszej kolejności montujemy klawiaturę matrycową, korzystając z długiego złącza GOLDPIN, które zapewnia jej niezbędne połączenie elektryczne oraz 4 tulejek dystansowych o wysokości 13 mm (wraz ze stosownymi śrubami/nakrętkami), które zapewnią jej konieczną stabilizację mechaniczną oraz pozycjonują ją w taki sposób, by znajdujące się na jej powierzchni klawisze umiejscowione były powyżej korpusu przełącznika.

Kolejnym krokiem procesu montażowego jest montaż modułu RFID/NFC, co wykonujemy w taki sposób, by znalazł się on najbliżej, jak to się da, od dolnej powierzchni montowanego w dalszym kroku wyświetlacza OLED (także korzystając z dedykowanego złącza GOLDPIN, tym razem krótkiego). W tym momencie ustawiamy w odpowiednich pozycjach przełączniki SMD umieszczone na module RFID, dzięki którym wybieramy aktywny interfejs komunikacyjny (w naszym wypadku SPI). Aby przygotować moduł do pracy z interfejsem SPI, przełączamy przełącznik oznaczony jako ON w pozycję On (blisko znacznika „1”), zaś przełącznik oznaczony jako KE w pozycję Off (blisko znacznika „KE”).

W ostatnim kroku montujemy wyświetlacz OLED, korzystając, jak wcześniej, z długiego złącza GOLDPIN zapewniającego mu niezbędne połączenie elektryczne oraz 4 tulejek dystansowych o wysokości 12 mm (wraz ze stosownymi śrubami/nakrętkami), które zapewniają mu konieczną stabilizację mechaniczną. Sam wyświetlacz przysłoni co prawda moduł RFID/NFC, lecz testy praktyczne wykazały, iż nie zawiera on elementów, które w sposób znaczący mogłyby zaburzyć funkcjonowanie anteny umieszczonej na tym peryferium.

Poprawnie zmontowany układ nie wymaga żadnych regulacji (poza koniecznością ustalenia kodu PIN/karty RFID/NFC administratora) i powinien działać tuż po włączeniu zasilania.

W przypadku zastosowania przekaznika innego typu niż przewidziany w projekcie, należy użyć tulejki dystansowej o innej wysokości, niż podano powyżej dostosowanej, co oczywiste, do wysokości korpusu wspomnianego elementu.

Na **fotografii 1** pokazano wygląd obwodu drukowanego urządzenia multiLock tuż przed przyłutowaniem modułu RFID, wyświetlacza OLED i klawiatury matrycowej (w wersji prototypowej, przed wprowadzeniem drobnych poprawek w zakresie położenia elementów).

Ustawienia Fuse-bitów:

CKSEL3...0: 0010
 SUT1...0: 10
 CKOUT: 1
 CKDIV8: 0
 EESAVE: 0

Obsługa

Konstruując algorytm obsługi urządzenia multiLock oraz wygląd graficznego interfejsu użytkownika, chciałem maksymalnie uprościć niezbędne czynności po stronie użytkownika systemu, zachowując jednocześnie spójność, jak i atrakcyjność wizualną poszczególnych ekranów funkcyjnych. Tuż po włączeniu zasilania urządzenie multiLock przechodzi do niezmiernie prostego ekranu głównego, w zakresie którego możemy wprowadzić PIN-kod użytkownika (lub administratora), jak też obsłużyć kartę RFID/NFC (odczytać jej numer seryjny). Co oczywiste, wprowadzenie poprawnego kodu PIN lub zbliżenie obsługiwanej karty RFID/NFC powoduje chwilowe załączenie wbudowanego przekaźnika, który w swoich zamierzeniach ma sterować rygłem elektromechanicznym.

Klawisze „0...9” służą w tym przypadku do wprowadzania kodu użytkownika, klawisz „*” kasuje wszystkie wprowadzone cyfry, zaś klawisz „#” wywołuje menu administratora, ale tylko wtedy, gdy aktywny jest tryb administratora. Wprowadzane cyfry kasowane są również automatycznie po czasie około 15 s bezczynności po stronie użytkownika.

Tryb administratora włączamy/wyłączamy (cyklicznie) z kolei za każdym razem wprowadzenia kodu administratora lub wczytania karty administratora, co sygnalizowane jest pojawieniem się ikonki klucza na ekranie graficznego interfejsu użytkownika. Po wejściu w tryb administratora dostępne mamy następujące opcje: dodawanie kodu/karty RFID/NFC użytkownika (wywoływane poprzez naciśnięcie przycisku „1”), usuwanie kodu/karty RFID/NFC użytkownika (wywoływane poprzez naciśnięcie przycisku „2”) oraz wyjście do ekranu głównego (wywoływane poprzez naciśnięcie przycisku „#”).

Po wejściu w menu dodawania kodu użytkownika ukazuje się prosty interfejs graficzny umożliwiający wprowadzenie nowo obsługiwanego kodu użytkownika czy też nowej karty RFID/NFC. Wprowadzanie i usuwanie cyfr przebiega analogicznie jak w przypadku ekranu głównego. Poprawne wykonanie operacji dodania kodu/karty użytkownika sygnalizowane jest dwukrotnym krótkim dźwiękiem wbudowanego buzzera piezoelektrycznego, zaś błędne wykonanie tejże operacji długim pojedynczym dźwiękiem wspomnianego buzzera.

Błąd w trakcie dodawania nowego kodu/karty RFID/NFC użytkownika może się pojawić wyłącznie w dwóch przypadkach: gdy nowo wprowadzany kod/karta został już

Tabela 1. Rodzaje błędów modułu PN532 i ich znaczenie funkcjonalne

Numer błędu	Rodzaj błędu
1	Błąd ramki potwierdzenia (ACK) modułu PN532
2	Przekroczenie czasu (1 s) odpowiedzi modułu PN532
3	Błąd nagłówka ramki wersji modułu PN532
4	Błąd konfiguracji RF
5	Błąd konfiguracji modułu SAM (Security Access Module)
6	Błąd odczytu tagu NFC
7	Błąd autoryzacji
8	Błąd odczytu bloku danych

wcześniej zapisany w pamięci urządzenia lub gdy wyczerpano limit miejsc w pamięci urządzenia przeznaczonych na tego rodzaju dane. Limit ten to odpowiednio (i niezależnie od siebie): 100 kodów PIN oraz 100 numerów seryjnych kart RFID/NFC. Wyjście z menu dodawania nowego kodu/karty użytkownika następuje poprzez naciśnięcie przycisku „#”.

Jeśli chodzi o drugą z opcji menu administratora, a mianowicie możliwość usuwania kodów/kart RFID/NFC użytkownika, to jest ona bliźniaczo podobna do dodawania tychże kart/kodów, przy czym w tym wypadku wystąpienie błędu towarzyszącego procesowi usuwania możliwie jest wyłącznie w przypadku, gdy usuwany kod/karta nie jest znany przez nasz system mikroprocesorowy, w związku z czym nie może zostać usunięty z pamięci.

Nie wspominałem jeszcze o specjalnym menu umożliwiającym wprowadzenie kodu/karty administratora, które wywoływane może być wyłącznie w trakcie włączania urządzenia, a następuje to na skutek założenia zworki ADMIN na stosownym goldpinie umieszczonym na obwodzie drukowanym. Po wejściu w menu dodawania kodu/karty RFID/NFC administratora postępujemy analogicznie, jak przy dodawaniu normalnego kodu PIN czy też karty użytkownika, przy czym w tym przypadku nie są generowane żadne sygnały błędów, gdyż wprowadzane dane nadpisują każdorazowo istniejące dane tego rodzaju. Wyjście z menu dodawania nowego kodu/karty administratora następuje poprzez naciśnięcie przycisku „#”.

Kolejną opcją dostępną w przypadku obecności zworki ADMIN na stosownym goldpinie umieszczonym na obwodzie drukowanym jest możliwość wykasowania całej pamięci kodów/numerów seryjnych kart, którą to opcję wywołujemy, wciskając i przytrzymując klawisz „*” podczas włączania urządzenia. W takim przypadku cała pamięć PIN-kodów i numerów seryjnych kart RFID/NFC (poza danymi autoryzacyjnymi administratora) zostanie wykasowana, na ekranie urządzenia pojawi się stosowna informacja, po czym system przejdzie do menu dodawania kodu/karty administratora, z którego, jak

wcześniej, wychodzimy, wciskając klawisz „#” (jeśli nie chcemy zmieniać danych autoryzacyjnych administratora).

Na koniec kilka słów uwagi należy się potencjalnej możliwości obsługi urządzeń NFC w rodzaju telefonu komórkowego. Jak wiadomo, niektóre z tych urządzeń wyposażono w interfejs NFC umożliwiający wymianę danych lub obsługę płatności zbliżeniowych. Czy nasz prosty system będzie w stanie odczytać numer seryjny takiego urządzenia? Oczywiście, że tak, bo musi ono spełniać standardy protokołu NFC. Problem jednak w tym, że systemy operacyjne telefonów komórkowych generują losowe numery seryjne (tzw. RID od Random ID), w związku z czym ich zapamiętywanie nie ma większego sensu, gdyż podczas kolejnego połączenia z takim urządzeniem wygeneruje ono zupełnie inny numer seryjny.

Zachowanie to można zmienić, ale przynajmniej w Androidzie nie jest to takie proste i wymaga modyfikacji firmware lub użycia specjalnych aplikacji. Ponoć zdarzają się telefony ze statycznym numerem UID (takie informacje odnalazłem na specjalistycznych forach), lecz ja na takie nie natrafiłem, więc należy założyć, że standardowo generują one losowe numery RID.

I na koniec jeszcze jedna, istotna informacja. Uruchamianiu urządzenia multiLock towarzyszy inicjalizacja wbudowanego modułu czytnika kart RFID/NFC. Jest to proces dość skomplikowany i jeśli zakończy się z jakiegoś powodu niepowodzeniem (np. nie ustawiliśmy na module odpowiedniego interfejsu komunikacyjnego), to stosowne informacje wraz z kodem błędu pokazane zostaną w ramach graficznego interfejsu użytkownika, zaś samo urządzenie przejdzie do trybu bezczynności (nie będzie można go użytkować). Obsługiwane kody błędów wraz z ich znaczeniem opisano w tabeli 1.

Na rysunku 3 pokazano z kolei kompletny diagram prezentujący sposób obsługi urządzenia multiLock i rodzaje dostępnych ekranów menu.

Robert Wołgajew, EP

[1] <https://ep.com.pl/projekty/projekty-ep/15235-nfc-lock>



Wygraj płytkę rozwojową Microchip PIC24F LCD i USB

Płytkę rozwojową Microchip PIC24F LCD i USB Curiosity Development Board (DM240018) to w pełni zintegrowana platforma rozwojowa, przeznaczona do odkrywania możliwości interfejsów z segmentowymi wyświetlaczami LCD oraz różnych funkcji mikrokontrolerów PIC24F. Układy te odznaczają się niskim poborem mocy, zintegrowanym interfejsem USB i zintegrowanym kontrolerem wyświetlaczy LCD.

Płytkę została starannie zaprojektowana od podstaw, aby w pełni wykorzystać możliwości zastosowanego mikrokontrolera – PIC24FJ512GU410 z kontrolerem USB i LCD. Zastosowane złącza i optymalne połączenia zapewniają dostęp do niezależnych urządzeń peryferyjnych (CIP). Wbudowane w układ scalony bloki CIP integrują wiele różnych funkcji systemowych w jednym MCU, upraszczając projekt i utrzymując niskie zużycie energii systemu oraz koszty BOM.

Kontroler LCD pozwala płynnie sterować wyświetlaczem LCD o aż 480 segmentach (64 segmenty i 8 linii wspólnych) nawet przy aktywnych trybach oszczędzania energii. Dzięki nowej funkcji – animacji niezależnej od rdzenia mikrokontrolera, na wyświetlaczu LCD można również wizualizować proste animacje, nawet jeśli MCU znajduje się w trybie oszczędzania energii. Płytkę pozwala na pomiar prądu urządzenia i pozwala

na uruchomienie zasilania z dodatkowej baterii pastylkowej.

Mikrokontroler PIC24FJ512GU410 jest 16-bitowym układem taktowanym z częstotliwością do 16 MHz. Oferuje do 512 kB pamięci Flash zorganizowanej w dwóch partycjach, obsługującej aktualizację OTA w czasie rzeczywistym i emulację pamięci EEPROM.

Na płytce znajduje się programator/debugger, zatem nie wymaga dodatkowego sprzętu do współpracy z MPLAB X IDE i MPLAB Code Configurator (MCC) firmy Microchip. Rozbudowane środowisko programistyczne oraz graficzne narzędzia konfiguracyjne obsługujące mikrokontrolery PIC24F pozwalają na przejście od koncepcji projektu do prototypu w bardzo krótkim czasie.

Aby mieć szansę na wygranie Microchip PIC24F LCD i USB Curiosity Development Board lub aby otrzymać kupon rabatowy



15% i bezpłatną wysyłkę, należy wypełnić formularz zgłoszeniowy na stronie: <https://page.microchip.com/E-Prak-LCD.html>.

Szczegółowe informacje na temat płytki rozwojowej można znaleźć na: <https://www.microchip.com/en-us/development-tool/dm240018>.

Natomiast dokumentacja mikrokontrolera jest dostępna na: <https://www.microchip.com/en-us/product/pic24fj512gu410>.



**Podstawowe parametry:**

- regulator poziomu zrealizowany na znanym i od wielu lat produkowanym układzie PGA2320 firmy Texas Instruments,
- aktywny układ regulacji tonów niskich i tonów wysokich z możliwością jego całkowitego ominięcia (bypass),
- selektor wejść z 3 wejściami stereofonicznymi przełączanymi miniaturowymi przełącznikami sygnałowymi,
- zbudowany na bazie wzmacniaczy operacyjnych zoptymalizowanych do zastosowań audio,
- zawiera przetwornik cyfrowo-analogowy na bazie stosunkowo taniego i bardzo dobrego układu PCM1794A.

* **Uwaga!** Elektroniczne zestawy do samodzielnego montażu. Wymagana umiejętność lutowania! Podstawową wersją zestawu jest wersja [B] nazywana potocznie KIT-em (z ang. zestaw). Zestaw w wersji [B] zawiera elementy elektroniczne (w tym [UK] – jeśli występuje w projekcie), które należy samodzielnie wlutować w dołączoną płytkę drukowaną (PCB). Wykaz elementów znajduje się w dokumentacji, która jest podlinkowana w opisie kitu. Mając na uwadze różne potrzeby naszych klientów, oferujemy dodatkowe wersje:

- **wersja [C]** – zmontowany, uruchomiony i przetestowany zestaw [B] (elementy wlutowane w płytkę PCB),
- **wersja [A]** – płytką drukowaną bez elementów i dokumentacji. Kity, w których występuje układ scalony wymagający zaprogramowania, mają następujące dodatkowe wersje:
- **wersja [A+]** – płytką drukowaną [A] + zaprogramowany układ [UK] i dokumentacja,
- **wersja [UK]** – zaprogramowany układ.

Dodatkowe materiały do pobrania ze strony www.ulubionykiosk.pl/media

- AVT6005-2 toneCtrl – regulator barwy dźwięku (2) (EP 10/2023)
- AVT6005-1 toneCtrl – regulator barwy dźwięku (1) (EP 9/2023)
- AVT5975 Regulator barwy dźwięku, głośności i balansu (EP 3/2023)
- AVT5873 Stereofoniczny aktywny regulator głośności (EP 8/2021)
- AVT5851 7-pasmowy korektor graficzny (EP 4/2021)
- AVT5816 Regulator balansu tonów (EP 10/2020)
- AVT5637 Wielokanałowy regulator głośności VCA (EP 8/2018)
- AVT5629 Cyfrowy regulator głośności z układem PT2257 (EP 6/2018)
- AVT3222 Sterowany dowolnym pilotem potencjometr audio z przełącznikiem (EdW 5/2018)

Nie każdy zestaw AVT występuje we wszystkich wersjach! Każda wersja ma załączony ten sam plik PDF! Podczas składania zamówienia upewnij się, którą wersję zamawiasz! <http://sklep.avt.pl>

W przypadku braku dostępności na stronie sklepu osoby zainteresowane zakupem płytek drukowanych (PCB) prosimy o kontakt via e-mail: kity@avt.pl.

Uniwersalny przedwzmacniacz (1)

Bardzo popularne wzmacniacze zintegrowane są zbudowane ze wzmacniaczy mocy i kompletnych układów przedwzmacniacza. Od długiego czasu są również oferowane oddzielne wzmacniacze mocy również w konfiguracji dual mono. Takie wzmacniacze mocy potrzebują do pracy osobnego przedwzmacniacza z układami selektora wejść, regulatoracji barwy i poziomu sygnału (siły głosu). Zaprezentowany projekt sprawdzi się doskonale w takiej konfiguracji.



Tor audio jest zbudowany z trzech podstawowych elementów: źródła sygnału, wzmacniacza i zestawów głośnikowych. Wśród fanów wysokiej jakości dźwięku chyba największe emocje wzbudzają wzmacniacze mocy i zestawy głośnikowe. Potem są źródła sygnału, głównie odtwarzacze CD i autonomiczne przetworniki cyfrowo-analogowe DAC. Entuzjaści płyt winylowych starają się stosować jak najlepsze wzmacniacze korekcyjne do wkładek dynamicznych.

Przedwzmacniacze są elementami toru, które w powszechnej opinii nie wpływają szczególnie na jakość sygnału i nie ma z nimi większych problemów. Wyjątkiem jest może potencjometr, a ściślej jego trwałość i współbieżność. Problemy ze zużyciem się tego elementu doprowadziły do zbudowania regulatorów z drabinki rezystorowej przełączanej wielostykowym przełącznikiem. Te „potencjometry” mają również wady w postaci

zużywających się styków. Żeby zapewnić w miarę płynną regulację, potrzebny jest przełącznik z dużą liczbą przełączników, co podnosi cenę i komplikuje układ.

Alternatywą jest scalony regulator. To rozwiązanie ma swoje zalety, ale też i wady. Zaletą jest duża trwałość, jednak takie elementy potrzebują do pracy sterownika mikroprocesorowego.

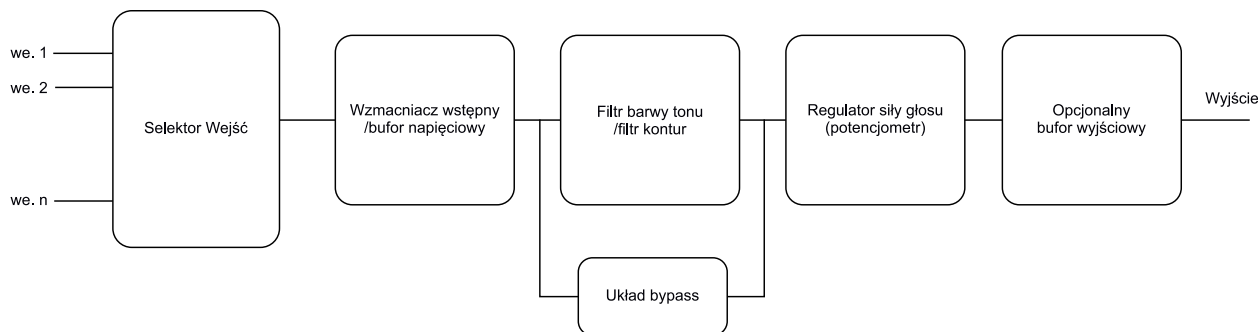
Budowa klasycznego przedwzmacniacza

Na **rysunku 1** pokazano schematycznie budowę klasycznego przedwzmacniacza audio. Układ ma do spełnienia kilka ważnych funkcji:

- dostosowanie poziomu sygnału wejściowego do czułości wejściowej wzmacniacza mocy,
- regulacja poziomu sygnału podawanego na wejście wzmacniacza mocy,

- selekcja źródeł sygnału wejściowego,
- modyfikacja częstotliwościowa sygnału: filtry kontur, filtry tonów niskich i wysokich itp.,
- dopasowanie impedancyjne pomiędzy źródłem sygnału wejściowego a poszczególnymi elementami przedwzmacniacza i wejściem wzmacniacza mocy.

Swoją nazwę przedwzmacniacz zawdzięcza realizacji pierwszej z wymienionych funkcji. Paradoksalnie teraz zazwyczaj nie ma potrzeby wstępnego wzmacniania poziomu sygnału źródłowego. Najczęściej stosowane źródła (odtwarzacze CD, autonomiczne przetworniki DAC i wzmacniacze korekcyjne) mają tak wysokie poziomy sygnału wyjściowego, a czułość wzmacniacza mocy jest tak duża, że w praktyce potrzebne jest tylko tłumienie sygnału. Tak na marginesie, brak konieczności stosowania



Rysunek 1. Schemat blokowy wzmacniacza audio

aktywnego stopnia wzmacniającego przyczynił się do powstania idei tak zwanego pasywnego przedwzmacniacza zawierającego tylko mechaniczny selektor wejść i klasyczny potencjometr. Z założenia miało to sprawić, że tor audio będzie lepszy, ponieważ brak jakiegokolwiek elektroniki to brak zniekształceń. To podejście w pewnych ściśle określonych sytuacjach może się obronić, ale w innych nie. Wróćmy do tego przy okazji omawiania dopasowania impedancyjnego.

Dostosowanie poziomu sygnału

Dostosowanie poziomu sygnału wejściowego nie musi oznaczać konieczności jego wzmocnienia, oczywiście jeżeli założymy, że przedwzmacniacz nie zawiera wzmacniacza korekcyjnego do gramofonu. Może istnieć konieczność wstępnego tłumienia dzielnikami rezystancyjnym zbyt dużego sygnału wejściowego tak, aby go dopasować do innych sygnałów w torze.

Zadaniem regulatora jest dostosowanie poziomu sygnału na wejściu wzmacniacza mocy do wymaganego natężenia dźwięku. Dwutorowy (stereofoniczny) regulator powinien się charakteryzować charakterystyką logarytmiczną i dobrą współbieżnością. Historycznie pierwszymi regulatorami były potencjometry obrotowe. Potem przyszła moda na potencjometry suwakowe, ale z powodu problemów z odpowiednim zabezpieczeniem ścieżek rezystancyjnych przed zabrudzeniem i kolejnej zmiany mody dzisiaj raczej się ich nie stosuje.

Dobry mechaniczny potencjometr obrotowy jest jednym z najlepszych wyborów w roli regulatora. Nie wprowadza szumów (prawie), nie wymaga zasilania i sterowania. Wadą jest ograniczona trwałość. Nawet najlepsze z nich czasem wymagają kłopotliwego serwisu (rozbierania, czyszczenia, smarowania), a po pewnym czasie pracy zużywają się ścieżki oporowe tak, że pozostaje tylko wymiana.

Remedium na te wady miały być przełączane drabinki rezystancyjne. Ale i w nich zużywa się przełącznik. Poza tym w tych prostszych rozwiązaniach zależnie od położenia regulatora zmienia się ich całkowita rezystancja, co jest na pewno niezbyt eleganckie

konstrukcyjnie, ale może również powodować problemy z dopasowaniem impedancji.

Kolejnym rozwiązaniem są specjalizowane scalone układy regulatorów. Te najlepsze mają bardzo dobre parametry, jeżeli chodzi o zniekształcenia nieliniowe i szumy. Charakteryzują się bardzo dużą niezawodnością, ale mają też wady. Szumia bardziej niż dobry potencjometr, wymagają zasilania i to często symetrycznego. Wymagają też sterownika z mikrokontrolerem. Komplikuje to układ, często też zasilanie powoduje wzrost ceny urządzenia.

Selektor sygnałów

Selektor sygnałów to najczęściej mechaniczny przełącznik obrotowy z wymaganą liczbą przełączanych sekcji. Po pewnym czasie pracy podobnie jak w klasycznym potencjometrze jego styki wymagają przeczyszczenia i nasmarowania. W lepszych rozwiązaniach stosuje się małosygnałowe przekaźniki. Ponieważ są hermetyczne, nie wymagają czynności serwisowych i pracują długo i bezawaryjnie. Spotykane są również cyfrowo sterowane analogowe klucze półprzewodnikowe.

Kształtowanie charakterystyki

Kolejnym elementem przedwzmacniacza są filtry. Kiedyś każdy szanujący się producent wzmacniaczy zintegrowanych umieszczał w torze filtry kształtujące charakterystykę częstotliwościową toru audio. Minimum stanowiły filtry końców pasma, czyli tonów niskich i wysokich oraz filtr konturu (*loudness*). Ten ostatni miał za zadanie korygować charakterystykę częstotliwościową toru zależnie od poziomu sygnału, tak by dostosować ją do charakterystyki częstotliwościowej naszego organu słuchu. We wzmacniaczach z lat 60. i początku 70. XX wieku stosowano też inne filtry, na przykład eliminujące zakłócenia mechaniczne generowane w gramofonach.

Od jakiegoś czasu w sprzęcie wysokiej klasy zaczęto unikać filtrów i traktować je jako źródło zniekształceń, głównie fazowych. Trzeba przyznać, że wiele takich układów nie było zbyt dobrze zaprojektowanych. Oferowały duże podbicia pasma i często nie dawały użytkownikowi możliwości

ich wyłączenia. Duże wzmocnienie tonów niskich i wysokich w połączeniu ze zbyt agresywnym konturem dawało koszmarny efekt, ale entuzjastów takiego grania nie brakowało.

Jednak usuwanie filtrów też nie jest dobrym rozwiązaniem. Szczególnie dotkliwie ich brak jest odczuwalny przy cichym słuchaniu muzyki, kiedy nasze uszy są mniej czułe na niskie i wysokie częstotliwości. Z wiekiem to zjawisko się tylko pogłębia. W dobrym przedwzmacniaczu mogą być dobrze zaprojektowane filtry z możliwością całkowitego ich wyłączenia (ominięcia). Użytkownik sam zdecydować, czy i kiedy ich użyje.

Dopasowanie impedancyjne

Dopasowanie impedancyjne polega na takim skonstruowaniu toru, żeby jego impedancja wejściowa była duża a wyjściowa mała. Duża impedancja wejściowa nie obciąża źródła sygnału, a mała impedancja wyjściowa pozwala bez problemów sterować wejściem wzmacniacza mocy. Możemy przyjąć, że duża impedancja to nie mniej niż 47 kΩ, a mała rezystancja to nie więcej niż 600 Ω. Technicznie nie ma żadnych problemów, żeby takie parametry osiągnąć w prosty sposób.

Uniwersalny przedwzmacniacz – założenia projektowe

Regulator poziomu sygnału

Jako regulator poziomu został wybrany znany i od wielu lat produkowany układ PGA2320 firmy Texas Instruments. Jest to stereofoniczny regulator głośności dźwięku przeznaczony do stosowania w sprzęcie profesjonalnym i konsumenckim. Konsekwencją tego wyboru jest konieczność zastosowania w układzie przedwzmacniacza mikroprocesorowego sterownika.

Regulator barwy

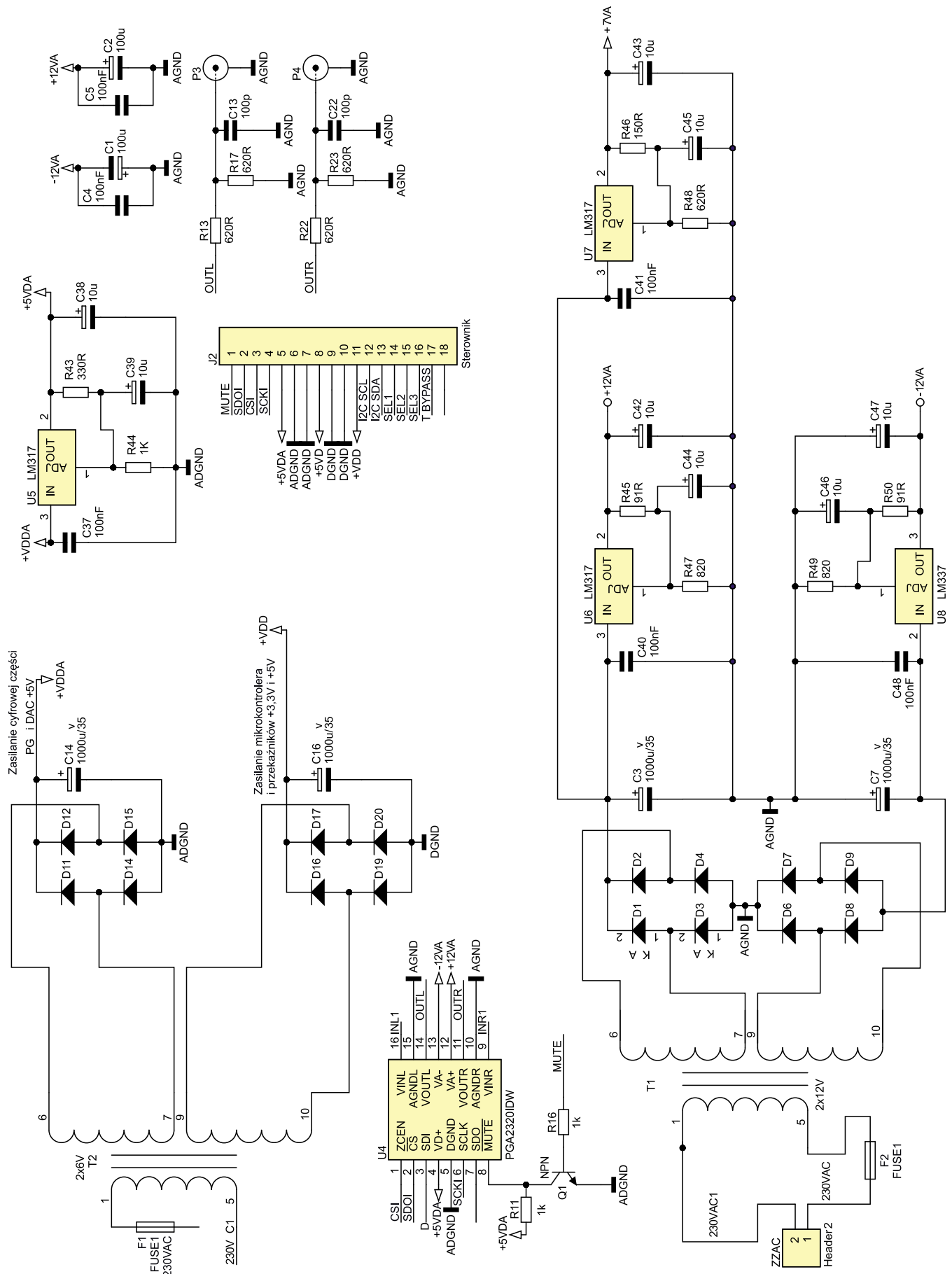
W założeniu przedwzmacniacz może być wyposażony w aktywny układ regulacji tonów niskich i tonów wysokich z możliwością jego całkowitego ominięcia (bypass). Układ filtrów powinien działać subtelnie i wprowadzać stałe i jak najmniejsze zniekształcenia fazowe w całym paśmie częstotliwości akustycznych. Regulacja jest realizowana za pomocą klasycznych potencjometrów stereofonicznych.

Selektor sygnału

Selektor wejść ma 3 wejścia stereofoniczne przełączane miniaturowymi przełącznikami

sygnałowymi i sterowane sterownikiem mikroprocesorowym. Dwa wejścia są przeznaczone dla analogowych zewnętrznych

sygnałów audio podłączanych do złącza Cinch, a jedno dla wewnętrznego źródła sygnału, jakim jest wyjście przetwornika DAC).



Rysunek 2. Schemat przedwzmacniacza z układem zasilania

Dopasowanie impedancyjne i fazowe

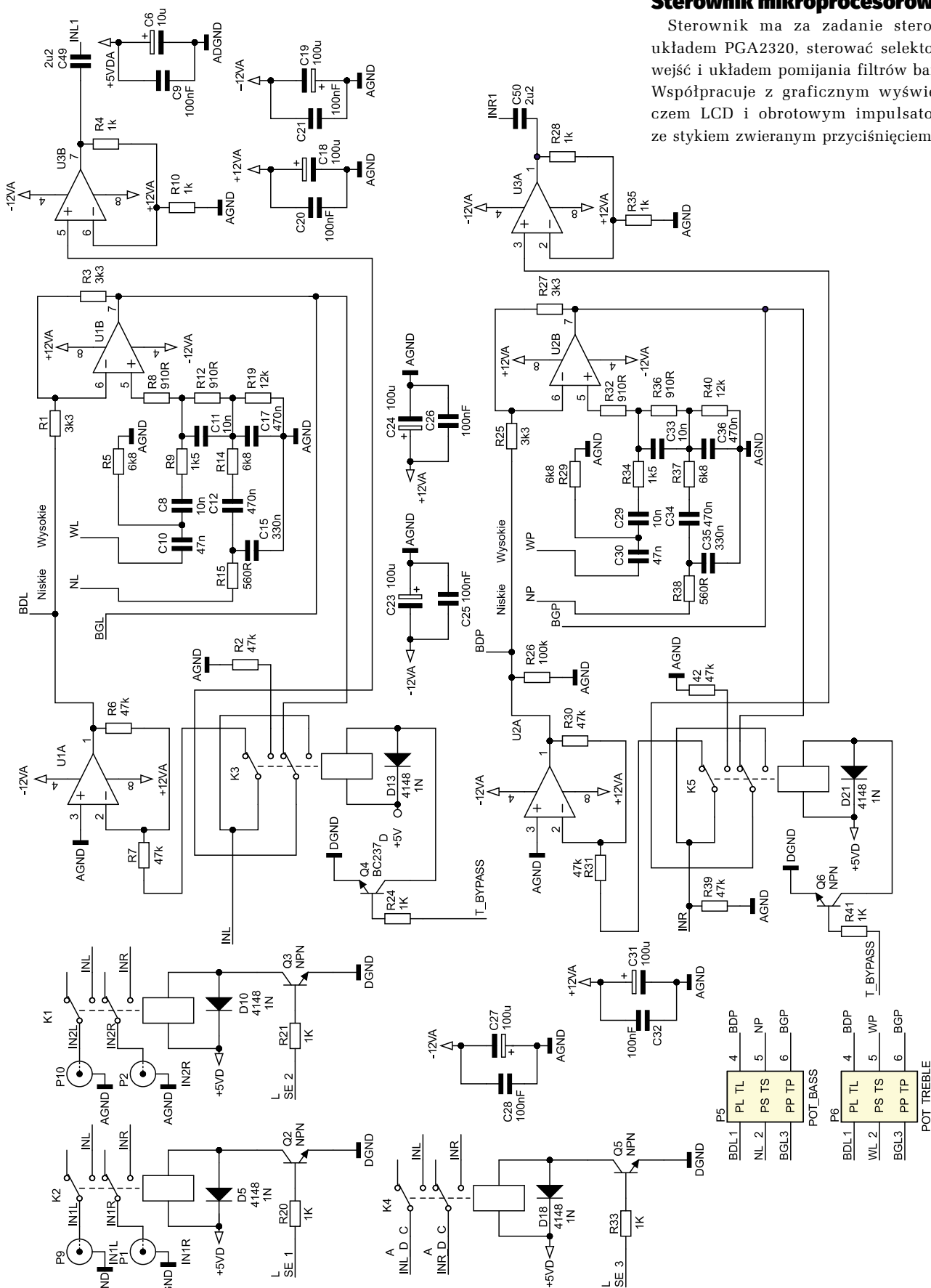
Tor audio, oprócz filtrów, zawiera układy wzmacniacza odwracają-

cego współpracującego z układem filtrów oraz wzmacniacza nieodwracającego zapewniającego małą impedancję wyjściową wymaganą do sterowania regulatorem PGA2320.

Wzmacniacze są zbudowane na bazie wzmacniaczy operacyjnych zoptymalizowanych do zastosowań audio.

Sterownik mikroprocesorowy

Sterownik ma za zadanie sterować układem PGA2320, sterować selektorem wejść i układem pomijania filtrów barwy. Współpracuje z graficznym wyświetlaczem LCD i obrotowym impulsatorem ze stykiem zwierzanym przyciśnięciem osi.



Rysunek 2. Schemat przedwzmacniacza z układem zasilania – cd.

Przewidziana jest układowa możliwość współpracy z odbiornikiem IR zdalnego sterowania. Oprogramowanie w pierwszej wersji ma sterować układem PGA2320, przełączać wejścia analogowe i sterować układem bypass barwy tonu. Wszystkie nastawy są zapisywane w nieulotnej pamięci EEPROM lub Flash i odtwarzane przy włączaniu urządzenia.

Przetwornik cyfrowo-analogowy

Przetwornik cyfrowo-analogowy nie jest klasycznym elementem przedwzmacniacza. Zazwyczaj jest wbudowany w odtwarzacz CD lub jest autonomicznym urządzeniem. Jednak umieszczenie DAC z wejściem USB lub SPDIF znacznie podnosi funkcjonalność, szczególnie przy odtwarzaniu materiału audio z komputera CD lub odtwarzacza CD. Technicznie możliwe jest odtwarzanie dobrej jakości streamingu przez łącze Bluetooth na przykład ze smartfonu.

Przyjąłem założenie, że w przedwzmacniaczu umieszczę układ przetwornika opartego o znany, stosunkowo tani i bardzo dobry przetwornik PCM1794 A. Źródłem sygnału może być moduł Amanero (wejście USB), konwerter Bluetooth/I²S, lub odbiornik SPDIF DIR9001.

Układ elektryczny

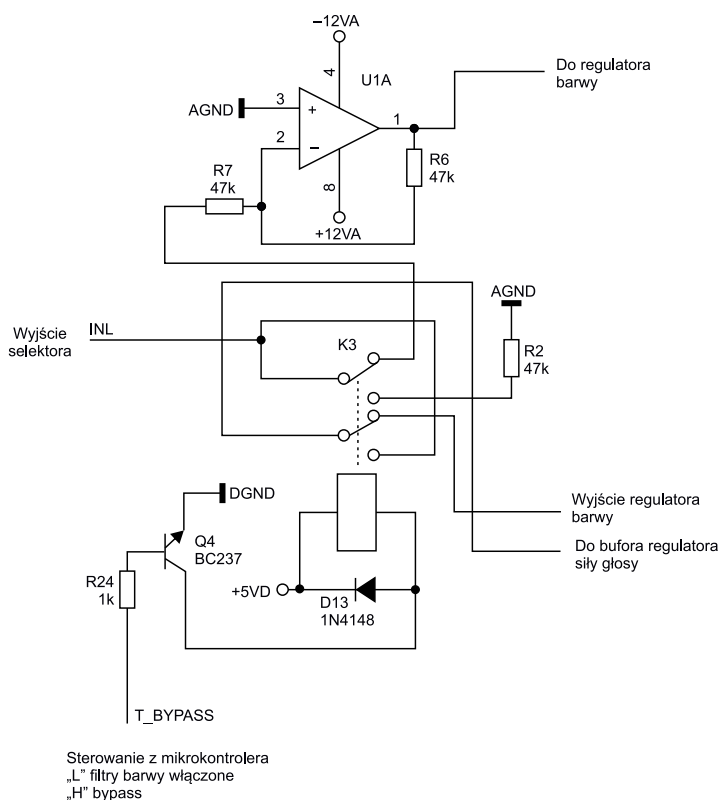
Układ elektryczny jest podzielony na 4 bloki:

- właściwy przedwzmacniacz z układami wzmacniaczy sygnału, regulatorem barwy tonu, regulatorem poziomu i selektorem wejść;
- układ zasilacza;
- kompletny przetwornik DAC z odbiornikiem S/PDIF i możliwością podłączenia konwertera USB/I²S lub Bluetooth/I²S;
- układ sterownika mikroprocesorowego.

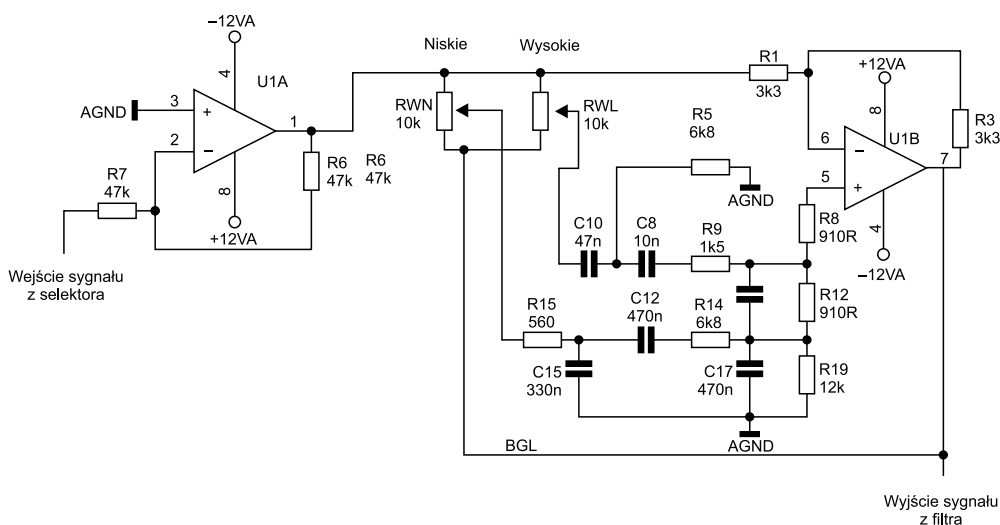
Dwa pierwsze bloki, czyli układy przedwzmacniacza i zasilacza, są umieszczone na jednej płycie, dlatego są pokazane na jednym schemacie ideowym – rysunek 2.

Prześledźmy drogę sygnału analogowego dla kanału lewego. Kanał prawy będzie identyczny.

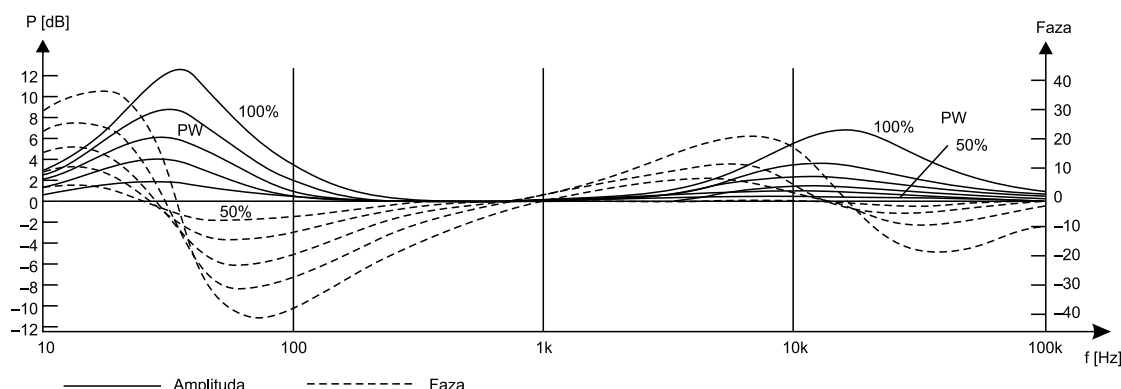
Sygnał z każdego z trzech stereofonicznych wejść IN1, IN2 lub OUT_DAC trafia do selektora wejść. Sygnały wejściowe z wejścia



Rysunek 3. Sterowanie wyłączaniem filtrów barwy (BYPASS)



Rysunek 4. Układ regulatora tonów



Rysunek 5. Charakterystyki amplitudowo-fazowe (źródło: artykuł „Radioelektronik” nr 8/1987)

IN1 są załączane przez przełącznik K2, z wejścia IN2 przez przełącznik K1, a z wyjścia DAC przez przełącznik K4. Selektorem steruje sterownik mikroprocesorowy. Wystawienie stanu wysokiego na jednym z wejść SEL1, SEL2 lub SEL3 powoduje przejście tranzystorów sterujących w stan nasycenia, zadziałanie przełącznika i podanie sygnału wejściowego do wyjścia selektora INL1 (kanał lewy) i punktu INR (kanał prawy). W danym momencie jest załączony tylko jeden przełącznik selektora. Dbą o to sterownik mikroprocesorowy.

Sygnał INL1 z wyjścia selektora jest podawany na styki przełącznika K3 pełniąc funkcję przełącznika sygnału. Podaje on na wejście bufora układu regulacji siły głosu sygnał bezpośrednio z wejść selektora lub poprzez układ regulacji barwy, realizując funkcję bypass – **rysunek 3**. Cewka przełącznika K3 jest sterowana z mikrokontrolera sygnałem T_BYPASS. Kiedy ten sygnał ma stan niski, tranzystor Q4 jest w stanie odcięcia i na cewkę przełącznika nie jest podawane napięcie +5 VD.

Styki przełącznika są ustawione w położeniu jak na rysunku 3. Górna para styków łączy sygnał INL z selektora z wejściem wzmacniacza odwracającego U1A (rezystor R7). Wyjście tego wzmacniacza jest połączone z wejściem filtra barwy tonu. Dolna para styków łączy wyjście regulatora barwy (niepokazanego na rysunku 3) z wejściem wzmacniacza bufora napięciowego sterującego regulatorem poziomu PGA2320. W takiej pozycji styków przełącznika K3 sygnał jest poddawany regulacji barwy tonu.

Kiedy sygnał T_BYPASS ma stan wysoki, tranzystor Q4 przechodzi w stan nasycenia i napięcie +5 VD zasila cewkę przełącznika. Styki przełącznika się przełączają. Górna para styków dołącza wejście INL przez rezystor R2 47 kΩ do masy, jednocześnie odłączając wejście wzmacniacza U1A od wejścia IN1L. Dolna para styków dołącza wejście IN1L do wejścia bufora sterującego PGA2320, jednocześnie odłączając wejście tego bufora od wyjścia regulatora barwy. Jak widać, tor regulatora jest tu całkowicie elektrycznie odcinany.

Wróćmy na chwilę do rezystora R2. Jego dołączenie do INL w stanie ominięcia regulatorów powoduje, że impedancja wejściowa ma ok. 47 kΩ. Taką samą impedancję wejściową ma wzmacniacz odwracający U1A włączony w tor regulatora barwy. Mamy więc w obu stanach przy aktywnym i nieaktywnym układzie bypass mniej więcej taką

samą impedancję wejściową 47 kΩ obciążającą źródła sygnału.

Schemat regulatora barwy został pokazany na **rysunku 4**. Jak już wiemy, sygnał jest najpierw podawany na wzmacniacz odwracający na układzie wzmacniacza operacyjnego U1A. Ten wzmacniacz ma do wykonania dwie funkcje. Zastosowany tu właściwy układ filtra na wzmacniaczu U1B pracuje poprawnie pod warunkiem, że źródło sygnału ma niską impedancję wejściową. Wzmacniacz U1A zapewnia spełnienie tego warunku. Druga funkcja to odwracanie fazy sygnału o 180°. Jest to niezbędne, ponieważ układ filtra bazuje na wzmacniaczu odwracającym. Sumarycznie cały tor filtra barwy tonu nie ma przesunięcia w fazie. Zapewnia to zgodność fazową toru, kiedy układ filtrów jest włączony i kiedy jest wyłączony. Oczywiście jeżeli pominiemy w rozważaniach nieuniknione przesunięcie fazowe samego filtra.

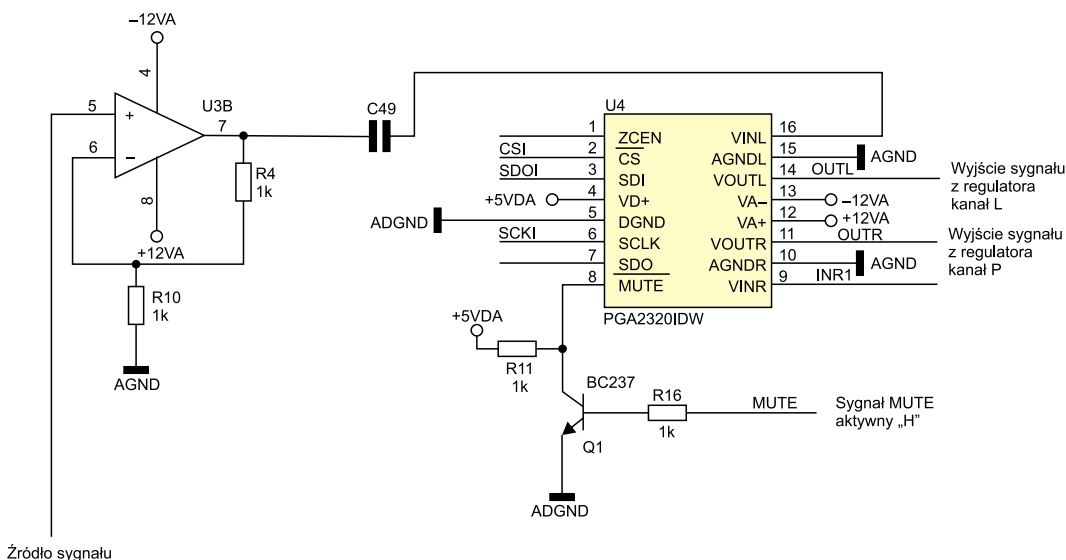
Sygnał z wyjścia U1A trafia na wejście aktywnego regulatora barwy z układem U1B. Większość układów regulacji barwy aktywnych i pasywnych wprowadza nieśtety większe lub mniejsze zniekształcenia. Najbardziej dokuczliwe jest przesunięcie fazowe wyraźnie zmieniające się w funkcji częstotliwości. Dość trudno jest znaleźć (nie mówiąc już o zaprojektowaniu) dobrze działający filtr barwy. Zbudowałem kilka takich filtrów i próbę czasu przetrwał tylko jeden z nich zastosowany tutaj. Jest kopią układu ze wzmacniacza Revox B-251 i został dokładnie opisany w artykule Marka Klimczaka z „Radioelektronika” numer 8 z 1987 roku.

Na **rysunku 5** pokazano charakterystyki amplitudowo-fazowe regulatora. Układ wprowadza wyjątkowo małe przesunięcia fazowe nieprzekraczające 40° dla 50 Hz. Regulacja tonów niskich wynosi ± 12 dB dla 30 Hz i $\pm 6,5$ dB dla 15 kHz. Jak widać na rysunku 5, sygnał po osiągnięciu maksymalnego podbicia dla tonów niskich zaczyna opadać, kiedy częstotliwość dalej

maleje. Podobnie jest z wysokimi częstotliwościami. Po osiągnięciu maksymalnego podbicia dla 15 kHz i przy dalszym wzmocnieniu częstotliwości amplituda zaczyna również opadać. Poza tym w zakresie częstotliwości od 150 Hz do 3 kHz sygnał amplitudowo i fazowo jest modyfikowany w minimalnym stopniu. Dlatego układ działa subtelnie i dźwięk jest bardzo dobrej jakości. Wyjście filtra nie powinno być obciążane wejściem o małej impedancji, ponieważ zmieni to charakterystyki amplitudowo-fazowe filtra.

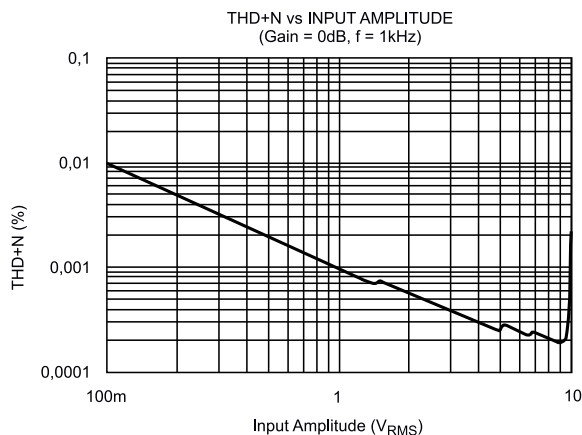
Sygnały z wejścia selektora (włączony bypass) lub z wyjścia regulatora barwy są podawane na wejście wtórnika napięciowego (bufora) zbudowanego na układzie U3B – **rysunek 6**. Komentarza wymaga topologia wzmacniacza z użyciem rezystorów R4 i R10. W takim układzie U3B pracuje jako wzmacniacz nieodwracający o wzmocnieniu ustalonym przez rezystory R4 i R10. Jeżeli nie włączymy rezystora R10, zewrzymy rezystor R4, to będziemy mieli układ wtórnika. Jednak umieszczenie tych rezystorów na płycie daje możliwość użycia U3B jako stopnia nieodwracającego i wzmacniającego o dużej impedancji wejściowej i małej wyjściowej. Dlaczego tak? Umożliwia to wstępne wzmocnienie sygnału przed PGA2320 i potem tłumienie go rezystorami R13 i R17 na wyjściu. Układ PGA2320 lepiej pracuje (ma mniejsze szumy i zniekształcenia) z dużymi sygnałami wejściowymi – **rysunek 7**. Katalogowo maksymalny poziom szumów PGA2311 wynosi 4 μ Vrms, a PGA2320 17,5 μ Vrms. Żeby uzyskać podobny stosunek sygnał/szum, musimy pracować z większymi amplitudami sygnału.

Mimo że użytkownicy zgłaszają w sieci, że PGA2320 wyraźnie szumi w porównaniu z PGA2311, ja tego nie zauważyłem w moim układzie, a ponieważ mam przedwzmacniacz z PGA2311 mogłem porównać obie konstrukcje. Docelowo wzmocniłem sygnał wejściowy 2× i na wyjściu stłumiłem go 2× rezystorami R13 i R17, oba po 620 Ω. To tłumienie



Źródło sygnału

Rysunek 6. Wzmacniacz napięciowy z regulatorem siły głosu



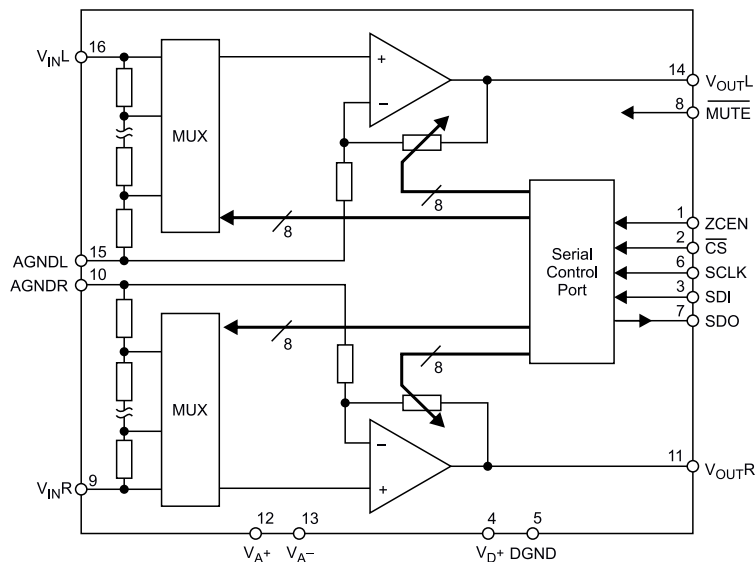
Rysunek 7. Zależność współczynnika zniekształceń w funkcji amplitudy wejściowej

można sobie dokładnie dobrać do czułości posiadanego wzmacniacza mocy, ale w praktyce nie jest to niezbędne.

Układ wtórnik (lub wzmacniacza nieodwracającego) zapewnia bardzo dużą impedancję wejściową, jak już wiemy, niezbędną do prawidłowej pracy regulatora barwy. W przypadku włączonego bypassu sygnał z selektora trafia bezpośrednio na ten wtórnik zapewniający również dużą impedancję wejściową obniżaną przez rezystor R2 do wartości 47 kΩ. Bufor ma również małą impedancję wyjściową konieczną do prawidłowej pracy regulatora siły głosu, czyli układu PGA2320. Według producenta źródło sygnału dołączone do wejścia PGA2320 powinno mieć impedancję nie mniejszą niż 600 Ω. Zbyt duża impedancja źródła również powoduje wzrost zniekształceń nieliniowych i wzrost szumów. Zastosowanie wtórnik (lub wzmacniacza nieodwracającego) powoduje, że warunek ten jest spełniony z dużym zapasem. Sygnał z wyjścia wtórnik trafia na wejście VINL PGA2320.

Na wejściu układu PGA2320 są włączone rezystory tworzące dzielnik tłumiący sygnał w zakresie od -95,5 dB do -0,5 dB. Rezystory są łączone przez macierz kluczy analogowych złączanych i wyłączanych przez wewnętrzne układy logiczne. PGA2320 nie jest tylko funkcjonalnym odpowiednikiem potencjometru. Oprócz tłumienia może również wzmacniać sygnał od +0,5 dB do +23,5 dB. Jeżeli skorzystamy z tej właściwości, to nie jest potrzebny dodatkowy układ wzmacniający, oczywiście jeżeli jest to konieczne. W praktyce PGA pracujący jako wzmacniacz może się stać przyczyną wzrostu poziomu szumów.

Schemat obwodu z PGA2320 został pokazany na **rysunku 8**. W torze sygnałowym PGA2320 jest umieszczony wzmacniacz operacyjny pracujący w konfiguracji wzmacniacza nieodwracającego. Wzmocnienie takiego układu wynosi $G=1+R2/R1$ (**rysunek 9**). Z tej zależności wynika, że wzmocnienie może osiągać wartość minimalną równą 1 dla $R2=0$ (wtórnik napięciowy). Stopień w tej konfiguracji może regulować przez zmianę rezystora R2 wzmocnienie



Rysunek 8. Schemat blokowy układu PGA2320

od 0 dB do +23,5 dB. W trakcie regulacji tłumienia od -95,5 dB do 0 dB w torze regulacji poziomu sygnału wzmacniacz pracuje jako bufor napięciowy (wzmocnienie 1, $R2=0$). Sygnał z wejścia jest tłumiony przez wejściowy dzielnik rezystancyjny – **rysunek 10**.

PGA2320 jest zasilany napięciem symetrycznym o maksymalnej wartości ± 15 V i napięciem +5 V (część cyfrowa). Możliwość zasilania napięciem symetrycznym o maksymalnej wartości ± 15 V była argumentem o użyciu tu PGA2320 zamiast PGA2311. Uprościło to układ zasilania, ponieważ nie trzeba było dodatkowych napięć ± 5 V niezbędnych do zasilania PGA2311. Sterowanie poziomem sygnału odbywa się poprzez magistralę SPI, a słowo danych ma długość 16 bitów.

Układ zasilania

Układ zasilania jest dość rozbudowany. Źródłem napięć przemiennych są dwa transformatory. Pierwszy transformator T1 dostarcza dwu symetrycznych napięć przemiennych o wartości +12 VAC przy

obciążeniu 500 mA przeznaczonych dla sekcji analogowej zasilacza. Ta sekcja zasilacza dostarcza napięć symetrycznych ± 12 V i napięcia +7 VA. Napięcia symetryczne są stosowane do zasilania wszystkich wzmacniaczy operacyjnych w układzie i części analogowej układu PGA2320.

Układ jest klasyczny, jego schemat został pokazany na **rysunku 11**. Napięcie przemienne jest prostowane w mostku Greatza i filtrowane kondensatorem elektrolitycznym 2200 μ F/25 V. Stabilizatorem napięcia dodatniego jest popularny stabilizator LT317 w wersji niskoszumnej, a napięcia ujemnego stabilizator LT337 również w wersji niskoszumnej. Tantalowy kondensator C44 (C46) o pojemności 10 μ F znacznie poprawia współczynnik tłumienia tętnień na wyjściu (*ripple rejection*). Według noty katalogowej LT317 A bez kondensatora ten współczynnik wynosi 65 dB, a po jego zastosowaniu wzrasta do typowej wartości 80 dB, co jest już bardzo dobrym wynikiem. Trochę gorzej jest z popularną wersją LT317, ponieważ tam ten

REKLAMA

Hurtownia elementów elektronicznych "AKSOTRONIK" zaprasza do swojego sklepu internetowego
Zaloguj się i kupuj ON-LINE na naszej stronie:

WWW.AKSOTRONIK.COM.PL

Aksotronik
ELEMENTY ELEKTRONICZNE

- Magnesy neodymowe oraz ferrytowe
Ceny od 0.10zł
- Przełączniki klawiszowe wodoszczelne, pyłoszczelne
Ceny od 2.40zł
- Druty oporowe od 0.16 do 0.81mm
Ceny od 5.70zł
- Prowadniki do przewodów
Ceny od 11.00zł
- Kostki elektryczne zaciskowe
Ceny od 0.22zł
- Szczotki węglowe do elektronarzędzi
Ceny od 2.60zł/kpl
- Przełączniki do elektronarzędzi zwykłe i elektromagnetyczne
Ceny od 7.00zł
- Złącza hermetyczne Supercol
Ceny od 1.10zł/kpl
- Pudełka/organizery
Ceny od 0.95zł
- Zestawy śrubek M2, M3 z nakrętkami i podkładkami
Ceny od 2.50zł

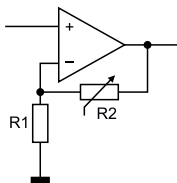
Uwaga!!! Powyższe ceny dotyczą zakupów minimalnych ilości hurtowych, poprzez nasz sklep internetowy.
W swojej ofercie posiadamy m.in.: półprzewodniki (diody, układy scalone, tranzystory, triaki, elementy optoelektroniczne), elementy dystansowe, złącza, przełączniki, elementy akustyczne, rezystory, kondensatory, kwarce, podstawki, moduły Arduino
Zapraszamy do kontaktu: **INFO@aksotronik.com.pl**, tel: (22) 783-20-51

współczynnik z kondensatorem ma średnią wartość 64 dB (nota katalogowa TI).

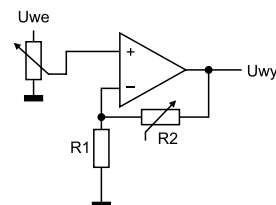
Napięcie wyjściowe jest ustalane rezystorami R45, R47 (R50, R49). Ja ustawiłem napięcia symetryczne o wartości ok. ± 12 V. Zmieniając dzielnik R45 i R47, można te napięcia zmienić, ustawiając na przykład napięcia ± 15 V. Może trzeba będzie zmienić transformator z wyższymi napięciami przemiennymi. W moim układzie po wyprostowaniu napięcie na kondensatorze C3 wynosi ok. 20 V. Układ U8 wymaga niewielkiego radiatora. Wyższy pobór prądu w ujemnej gałęzi wynika z poboru konwertera U/I przetwornika analogowo-cyfrowego.

Z dodatniego bieguna kondensatora C3 jest pobierane napięcie wejściowe stabilizatora napięcia ok. +7 V zbudowanego z układem U7 typu LM317. To napięcie jest stosowane przez układ zasilania przetwornika do wytwarzania napięć +5 V zasilających obwody analogowe przetwornika PCM1794 A. Powiemy o tym dokładniej przy okazji omawiania układu przetwornika. Sekcja analogowa zasilacza ma swoją masę oznaczoną AGND.

Drugi transformator T2 dostarcza dwu napięć przeznaczonych do zasilania układów cyfrowych. Pierwsze z uzwojeń przeznaczone jest do zasilania układów cyfrowych układu PGA2320 i przetwornika PCM1794 A. Napięcie to po wyprostowaniu i odfiltrowaniu kondensatorem C14 jest podawane na wejście stabilizatora +5 V U5 (LM317). Napięcie wyjściowe +5 VDA zasila obwody cyfrowe układu PGA2320. Napięcie +5 VDA jest również napięciem wejściowym dla stabilizatora



Rysunek 9. Układ wzmacniacza nieodwracającego



Rysunek 10. Regulacja poziomu sygnału w PGA2320

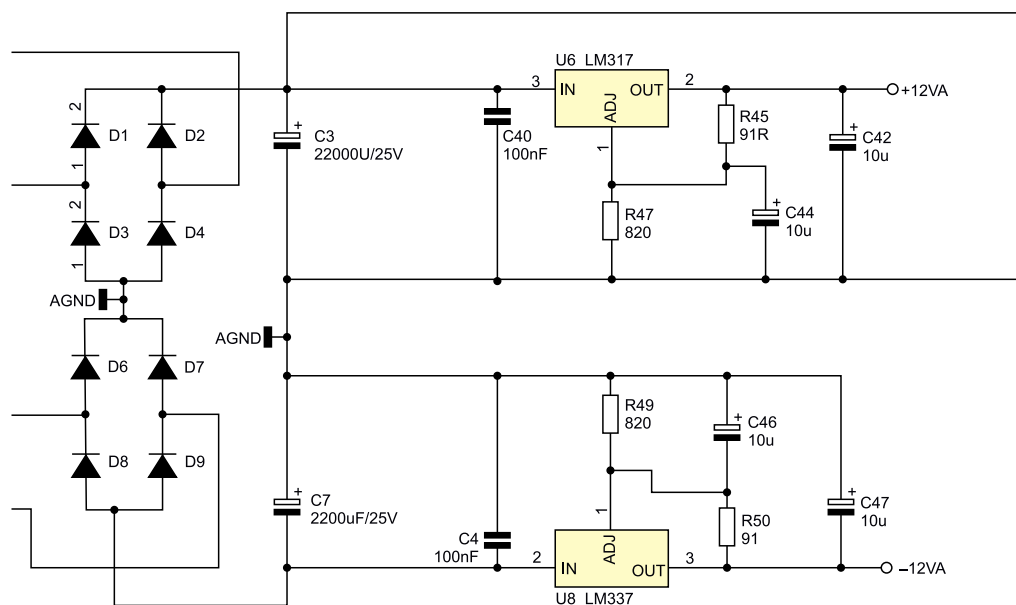
+3,3 V zasilającego układy cyfrowe przetwornika PCM1794 A.

Drugie z uzwojeń jest przeznaczone do zasilania układów cyfrowych sterownika mikroprocesorowego. Jest prostowane w mostku i filtrowane kondensatorem C16 1000 μ F/25 V. Stabilizatory napięć zasilających sterownik są już w układzie sterownika. To napięcie ma swoją masę DGND galwanicznie izolowaną od mas

układów analogowych AGND i DAGND. Zapobiega to przedostawaniu się zakłóceń z układu sterownika do obwodów analogowych przedwzmacniacza.

Na tym etapie kończymy pierwszą część opisu tego interesującego projektu. W kolejnej części omówimy budowę bloku przetwornika PCM1794A oraz sterownika mikroprocesorowego.

Tomasz Jabłoński, EP



Rysunek 11. Stabilizator napięć symetrycznych ± 12 V

REKLAMA

Świat projektantów i programistów
dla elektroniki w nowej odsłonie.
Odwiedź wiecznie młody

ELPORTAL.pl

**Podstawowe parametry:**

- sygnalizacja przerwy w obwodzie – rezystancji powyżej 2 kΩ pomiędzy zaciskami pomiarowymi,
- brak przepływu prądu oznajmiany przez głośne pischczenie i świecenie czerwonej diody LED,
- przepływ prądu między zaciskami wywołuje świecenie zielonej diody LED,
- zasilanie napięciem stałym 3 V z pojedynczej baterii CR2032,
- wbudowany wyłącznik zasilania,
- pobór prądu 5...20 mA.

* **Uwaga!** Elektroniczne zestawy do samodzielnego montażu. Wymagana umiejętność lutowania! Podstawową wersją zestawu jest wersja [B] nazywana potocznie KIT-em (z ang. zestaw). Zestaw w wersji [B] zawiera elementy elektroniczne (w tym [UK] – jeśli występuje w projekcie), które należy samodzielnie wylutować w dołączoną płytkę drukowaną (PCB). Wykaz elementów znajduje się w dokumentacji, która jest podlinkowana w opisie kitu. Mając na uwadze różne potrzeby naszych klientów, oferujemy dodatkowe wersje:

- wersja [C] – zmontowany, uruchomiony i przetestowany zestaw [B] (elementy wylutowane w płytkę PCB),
 - wersja [A] – płytka drukowana bez elementów i dokumentacji.
- Kity, w których występuje układ scalony wymagający zaprogramowania, mają następujące dodatkowe wersje:
- wersja [A+] – płytka drukowana [A] + zaprogramowany układ [UK] i dokumentacja,
 - wersja [UK] – zaprogramowany układ.

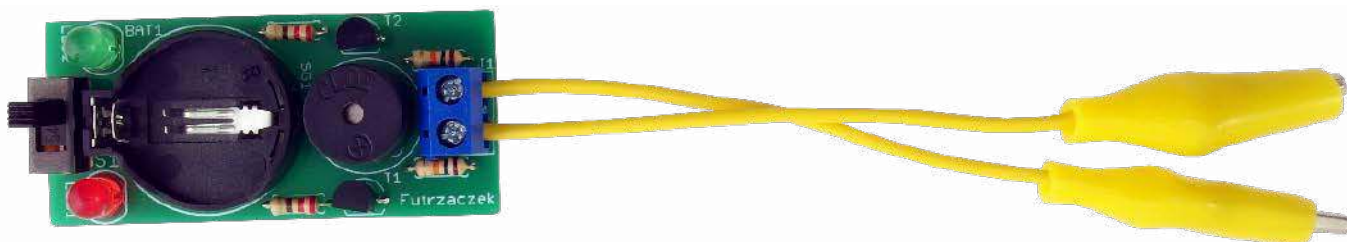
Dodatkowe materiały do pobrania ze strony www.ulubionykiosk.pl/media

- AVT1862 Sygnalizator akustyczny (EP 8/2015)
 AVT1450 Sygnalizator (nie)włączonych świateł w samochodzie (EP 5/2007)

Nie każdy zestaw AVT występuje we wszystkich wersjach! Każda wersja ma załączony ten sam plik PDF! Podczas składania zamówienia upewnij się, którą wersję zamawiasz!
<http://sklep.avt.pl>

W przypadku braku dostępności na stronie sklepu osoby zainteresowane zakupem płytek drukowanych (PCB) prosimy o kontakt via e-mail: kity@avt.pl.

W ofercie AVT*

AVT6011

Sygnalizator utraty ciągłości

Większość multimetrów ma wbudowany „pischczyk”, którym można szybko skontrolować, czy przez badany obwód może płynąć prąd. Zaprezentowany tu prosty układ robi coś zupełnie przeciwnego – zaczyna sygnalizować światłem i dźwiękiem, że ciągłość obwodu została przerwana. Jest lekki, kompaktowy i prosty w obsłudze – przyda się każdemu!

Fakt przepływu prądu przez jakiś przewód lub ścieżkę na płytce drukowanej można stwierdzić bardzo łatwo – większość multimetrów ma odpowiedni sygnalizator, można też na piechotę posłużyć się baterijką i żarówką. Multimetr pischczy (lub żarówka świeci), czyli prąd płynie. Co zrobić w sytuacji, kiedy przez jakiś obwód może płynąć prąd elektryczny, ale nam zależy na informacji

o przerwaniu jego ciągłości? Ciągłe wsłuchiwanie się w pischczenie albo patrzenie na żarówkę może być irytujące. Opisany układ zawiadomi użytkownika, kiedy obwód zostanie otwarty.

Gdzie to może się przydać? Na przykład podczas napraw połączeń przewodów w puszkach elektrycznych lub przy sprawdzaniu, czy przewód nie jest złamany. Jeżeli prąd

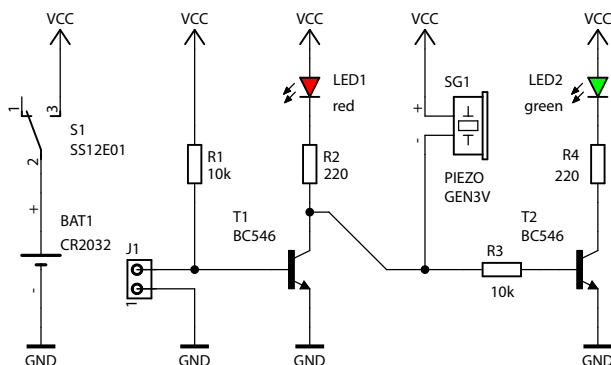
nie będzie mógł płynąć, układ poinformuje o tym zarówno światłem, jak i dźwiękiem.

Budowa i działanie

Schemat ideowy omawianego układu znajduje się na **rysunku 1**. Zasilanie stanowi pojedyncza bateria typu CR2032, mocowana w odpowiednim koszyku, która jest odcinana od reszty obwodu za pomocą niewielkiego przełącznika S1. Po rozwarciu jego styków układ nie pobiera jakiegokolwiek prądu z tej baterii.

Badany obwód podłącza się do zacisków złącza J1. Jeżeli jego rezystancja jest bardzo wysoka, wówczas prąd płynący przez rezystor R1 do bazy tranzystora T1 powoduje nasycenie tego elementu półprzewodnikowego. To z kolei skutkuje załączeniem diody LED1 (czerwonej) i sygnalizatora SG1. Ponieważ napięcie kolektor-emiter T1 staje się bardzo małe, rzędu kilkudziesięciu miliwoltów, tranzystor T2 zostaje zatkany – jego napięcie baza-emiter jest zbyt małe, by go otworzyć. Dlatego dioda LED2 nie świeci.

Po zwarceniu zacisków złącza J1, czyli w pożądanym przypadku, rolę tranzystorów odwracają się. Podłączona do J1 niewielka rezystancja tworzy z R1 dzielnik napięcia, który jest zbyt niski do otwarcia tranzystora T1. W tej sytuacji potencjał jego kolektora jest zbliżony



Rysunek 1. Schemat ideowy sygnalizatora utraty ciągłości

Wykaz elementów: (kupuj na stronie sklep.avt.pl lub osobiście Warszawa, ul. Leszczyńska 11, tel. +48222578451, e-mail: handlowy@avt.pl)

Rezystory: (THT o mocy 0,25 W)
 R1, R3: 10 kΩ
 R2, R4: 220 Ω

Półprzewodniki:

LED1: dioda LED czerwona 5 mm
 LED2: dioda LED zielona 5 mm

T1, T2: BC546

Pozostałe:

BAT1: koszyk CR2032 THT poziomy (np. KOSZYK BAT 6) +
 bateria CR2032
 J1: ARK2/500
 S1: SS12E01

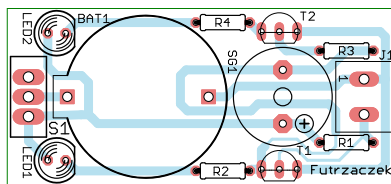
SG1: PIEZO GEN3V
 Dwa odcinki przewodu LgY 0,50 mm² np. LGY0.50
 ŻOŁTY (opis w tekście)
 Dwa krokodyłki izolowane np. KROK MET IZ3 (opis w tekście)

do 3 V. Niewielki prąd bazy, jaki przepływa przez LED1 i SG1 (głównie przez sygnalizator) do bazy T2, powoduje zaświecenie diody LED2. Jednocześnie natężenie prądu bazy T2 jest na tyle niskie, że LED1 i SG1 nie działają.

Montaż i uruchomienie

Układ został zmontowany na jednostronnej płytce drukowanej o wymiarach 25×55 mm. Jej schemat został pokazany na **rysunku 2**. Nie ma na niej otworów montażowych, układ można zaizolować, na przykład rurką termokurczliwą o dużym przekroju.

Montaż układu jest bardzo prosty i nawet początkującym elektronikom nie zajmie wiele



Rysunek 2. Schemat płytki PCB

czasu. Polecam rozpocząć od elementów najniższych, czyli rezystorów.

Poprawnie zmontowany układ jest gotowy do działania po włożeniu baterii CR2032 do koszyka na płytce. Pobór prądu przez układ wynosi około 20 mA przy rozwarciach zaciskach złącza J1 oraz około 5 mA przy zwartych zaciskach. Jako graniczną wartość

rezystancji, przy której układ rozpoznaje obwód jako zwarty, przyjęto 2 kΩ. Powyżej tej granicy tranzystor T1 nie nasycy się, co prowadzi do świecenia dwóch diod jednocześnie, co też można zastosować jako swego rodzaju informację o stanie badanego obwodu.

W układzie prototypowym do złącza J1 zostały podłączone krótkie (około 10 cm) przewody zakończone krokodylkami. Takie rozwiązanie znacząco ułatwia dołączanie tego układu do monitorowanych przewodów lub styków. Wygląd całego urządzenia pokazuje fotografia tytułowa, choć to tylko propozycja – do zacisków złącza J1 można dołączyć mnóstwo innych końcówek.

Michał Kurzela, EP



Podstawowe parametry:

- bazuje na układzie PAM8320,
- jest wyposażony w typowe zabezpieczenia chroniące przed uszkodzeniem,
- dysponuje mocą do 20 W/4 Ω przy typowym zasilaniu 12 V,
- zasilanie z zakresu 4,5...15 V,
- w standardowych zastosowaniach nie wymaga użycia radiatora.

*** Uwaga!** Elektroniczne zestawy do samodzielnego montażu. Wymagana umiejętność lutowania! Podstawową wersją zestawu jest wersja [B] nazywana potocznie KIT-em (z ang. zestaw). Zestaw w wersji [B] zawiera elementy elektroniczne (w tym [UK] – jeśli występuje w projekcie), które należy samodzielnie wstawić w dołączoną płytkę drukowaną (PCB). Wykaz elementów znajduje się w dokumentacji, która jest podlinkowana w opisie kitu. Mając na uwadze różne potrzeby naszych klientów, oferujemy dodatkowe wersje:

- wersja [C] – zmontowany, uruchomiony i przetestowany zestaw [B] (elementy wstawiane w płytkę PCB),
 - wersja [A] – płyta drukowana bez elementów i dokumentacji.
- Kity, w których występuje układ scalony wymagają zaprogramowania, mają następujące dodatkowe wersje:
- wersja [A+] – płyta drukowana [A] + zaprogramowany układ [UK] i dokumentacja,
 - wersja [UK] – zaprogramowany układ.

Dodatkowe materiały do pobrania ze strony www.ulubionykiosk.pl/media

- AVT5922 Wzmacniacz audio dla wymagających części 1 i 2 (EP 3/2022)
- AVT5836 Cyfrowy wzmacniacz mocy stereo z interfejsem I²S (EP 1/2021)
- AVT5756 Cyfrowy wzmacniacz mocy z interfejsem Bluetooth (EP 4/2019)
- AVT5717 Opóźniacz dołączenia głośników zasilany 230 V (EP 9/2019)
- AVT5669 Wzmacniacz mocy audio 4×48 W/4 Ω (EP 4/2019)
- AVT1982 Wzmacniacz z kanałem basowym 2.1 (EP 1/2019)
- AVT1973 Uniwersalny, stereofoniczny wzmacniacz mocy 2×10 W/8 V z regulacją barwy dźwięku (EP 2/2018)
- Miniatury, stereofoniczny wzmacniacz mocy (EP 10/2017)

Nie każdy zestaw AVT występuje we wszystkich wersjach! Każda wersja ma załączony ten sam plik PDF! Podczas składania zamówienia upewnij się, którą wersję zamawiasz! <http://sklep.avt.pl>

W przypadku braku dostępności na stronie sklepu osoby zainteresowane zakupem płytek drukowanych (PCB) prosimy o kontakt via e-mail: kity@avt.pl.

Mikrowzmacniacz mocy 20 W na układzie PAM8320

Zaprezentowany minimoduł jest kompletnym wzmacniaczem średniej mocy audio pracującym w klasie D. Wzmacniacz może posłużyć do realizacji mobilnych systemów nagłośnieniowych, może pracować samodzielnie np. jako wzmacniacz do systemu PC-audio lub jako zamiennik uszkodzonych i niedostępnych wzmacniaczy niewielkiej mocy w serwisowanych sprzętach RTV.

Układ PAM8320 jest monofoniczną końcówką mocy pracującą w konfiguracji mostkowej, wyposażoną w zabezpieczenia chroniące przed uszkodzeniem. PAM8320 pracuje poprawnie w szerokim zakresie napięć zasilania 4,5...15 V, współpracując z obciążeniem 4 Ω. W zależności od obciążenia i zasilania dysponuje mocą do 20 W/4 Ω przy typowym zasilaniu 12 V. W przypadku wzmacniania sygnałów muzycznych, ze względu na wysoką sprawność, układ nie wymaga stosowania

radiatora, do odprowadzania ciepła – wystarczy miedź płytki drukowanej. Budowę wewnętrzną układu pokazano na **rysunku 1**.

Budowa i działanie

Schemat układu wzmacniacza pokazano na **rysunku 2**. Aplikacja nie odbiega od noty katalogowej Diodes Inc. Sygnał wejściowy doprowadzony jest do złącza IN, rezystor R4 umożliwia korekcję wzmocnienia układu (w modelu wynosi 15 dB), kondensator



C13 separuje składową stałą i określa dolne pasmo przenoszenia.

Wyjścia mostka doprowadzone są poprzez filtr FB1, FB2/C9, 10 do zacisków wyjściowych

Wykaz elementów: (kupuj na stronie sklep.avt.pl lub osobiście Warszawa, ul. Leszczyńska 11, tel. +48222578451, e-mail: handlowy@avt.pl)

Rezystory:

R1: 10 Ω 0,25 W (SMD1206)
R2, R3, R4: 10 kΩ 1% (SMD0603)

Kondensatory:

C1, C5, C6: 10 μF/25 V ceramiczny (SMD1206)
CE1, CE2: 470 μF elektrolityczny Low ESR (CED10.0P5.0)

C2, C3, C4: 0,1 μF/25 V ceramiczny (SMD0603)

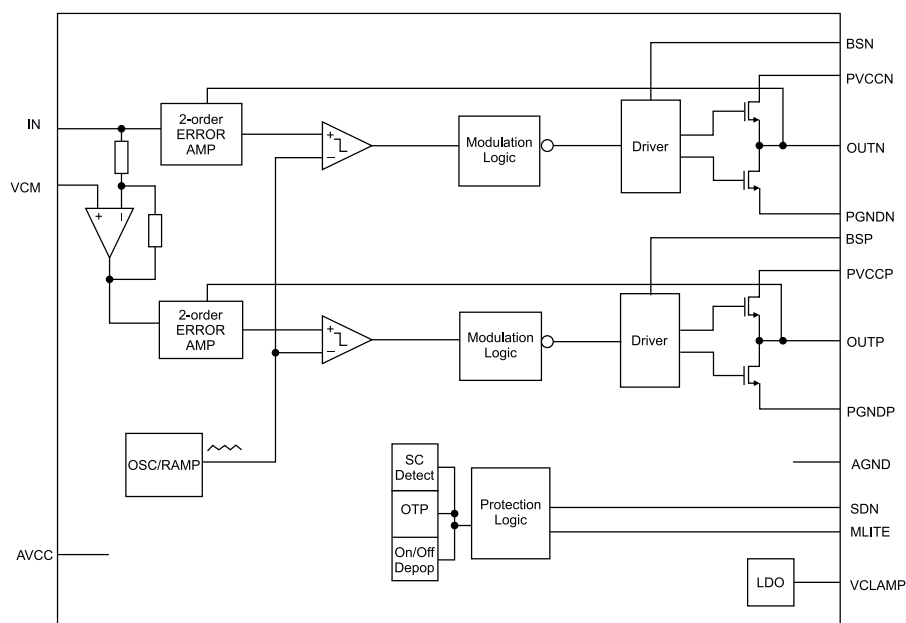
C7, C8, C11, C12, C13: 1 μF/25 V, ceramiczny (SMD0603)
C9, C10: 10 nF/25 V ceramiczny (SMD0603)

Półprzewodniki:

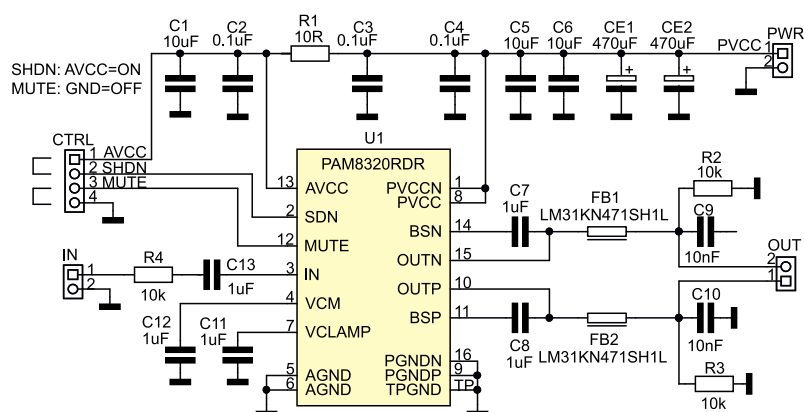
U1: PAM8320RDR (SO-16EP)

Pozostałe:

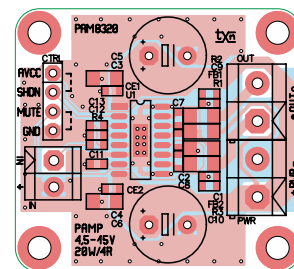
CTRL: złącze SIP4
FB1, FB2: dławik ferrytowy BLM31KN471S1L (SMD1206)
IN: CONN złącze śrubowe DG 2 piny, 3,5 mm (DG381-3.5-2)
OUT, PWR: złącze śrubowe DG 2 piny, 5,0 mm (DG126-5.0-2)



Rysunek 1. Budowa wewnętrzna PAM8320



Rysunek 2. Schemat ideowy wzmacniacza



Rysunek 3. Schemat płytki PCB

OUT. PAM8320 ma wyprowadzenie sterujące trybem obniżonego poboru mocy SDN, aktywowane stanem niskim oraz wyprowadzenie wyciszania MUTE, aktywowane stanem wysokim. Oba sygnały sterują układem bez dodatkowych zakłóceń podczas wyciszania lub zmiany trybu obniżonego poboru mocy. Sygnały SHDN/MUTE wraz z zasilaniem AVCC doprowadzone są do złącza CTRL i umożliwiają sterowanie trybem pracy końcówki mocy.

Zasilanie układu PVCC z zakresu 4,5...15 V (typowo 12 V) doprowadzone jest przez złącze PWR i filtrowane poprzez CE1/2, C1...4.

Montaż i uruchomienie

Wzmacniacz zmontowany jest na miniaturowej dwustronnej płytce drukowanej, której schemat został pokazany na **rysunku 3**. Układ zmontowany ze sprawdzonych elementów działa po włączeniu zasilania i podaniu stanu wysokiego na wyprowadzenie SHDN złącza CTRL oraz stanu niskiego na wyprowadzenie MUTE (np. po zwarceniu zwoją AVCC z SHDN i MUTE z GND). Przyjemnego odsłuchu...

Adam Tatuś, EP

REKLAMA



Elektronika
Praktyczna

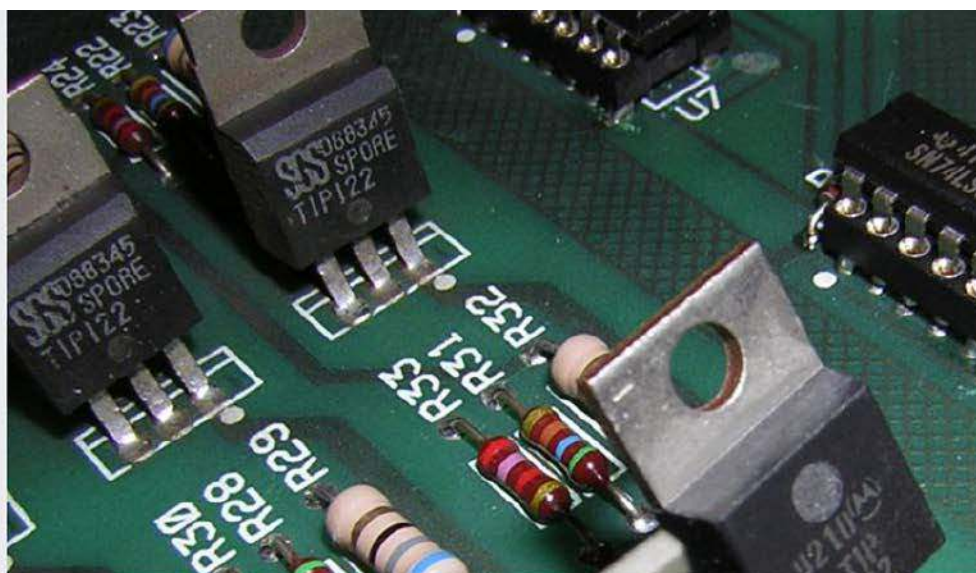
@ElektronikaPraktyczna

Strona główna

Posty

Filmy

Zdjęcia



Lubię to! Udostępnij Zaproponuj zmiany

www.facebook.com/ElektronikaPraktyczna

**Podstawowe parametry:**

- cztery wyjścia typu HighSide 12 V/5 A,
- zastosowano specjalizowane klucze półprzewodnikowe HighSide typu BTS5030,
- napięcie obciążenia powinno mieścić się w zakresie 6...15 V,
- dwa wyjścia mają możliwość pomiaru prądu obciążenia.

*** Uwaga!** Elektroniczne zestawy do samodzielnego montażu. Wymagana umiejętność lutowania! Podstawową wersją zestawu jest wersja [B] nazywana potocznie KIT-em (z ang. zestaw). Zestaw w wersji [B] zawiera elementy elektroniczne (w tym [UK] – jeśli występuje w projekcie), które należy samodzielnie wlutować w dołączoną płytkę drukowaną (PCB). Wykaz elementów znajduje się w dokumentacji, która jest podlinkowana w opisie kitu. Mając na uwadze różne potrzeby naszych klientów, oferujemy dodatkowe wersje:

Dodatkowe materiały do pobrania ze strony www.ulubionykiosk.pl/media

- Filtr zasilania dla Raspberry Pi (EP 9/2023)
- Ekspander GPIO RPi z taśmą FPC (EP 8/2023)
- Sterownik unipolarnego mikrosilnika krokowego dla Pi Pico (EP 7/2023)
- Sterownik dwóch silników krokowych do Raspberry Pi (EP 6/2023)
- Expander wyjść z PWM na bazie układu PCA9624 (EP 6/2023)
- Moduł redundancji zasilania dla Raspberry Pi Zero (EP 5/2023)
- Sterownik dwóch mikrosilników krokowych do Pi Zero (EP 3/2023)
- Sterownik taśm LED RGBCCCT 12 V dla RPi Zero (EP 3/2023)
- Eliminatory drgań styków mechanicznych (EP 1/2023)
- Moduł redundancji zasilania do komputerów SBC (EP 1/2023)
- Sterownik mikrosilnika krokowego dla Pi Pico (EP 12/2022)

- wersja [C] – zmontowany, uruchomiony i przetestowany zestaw [B] (elementy wlutowane w płytkę PCB),
 - wersja [A] – płytka drukowana bez elementów i dokumentacji.
- Kity, w których występuje układ scalony wymagający zaprogramowania, mają następujące dodatkowe wersje:
- wersja [A+] – płytka drukowana [A] + zaprogramowany układ [UK] i dokumentacja,
 - wersja [UK] – zaprogramowany układ.

Nie każdy zestaw AVT występuje we wszystkich wersjach! Każda wersja ma załączony ten sam plik PDF! Podczas składania zamówienia upewnij się, którą wersję zamawiasz! <http://sklep.avt.pl>

W przypadku braku dostępności na stronie sklepu osoby zainteresowane zakupem płytek drukowanych (PCB) prosimy o kontakt via e-mail: kity@avt.pl.

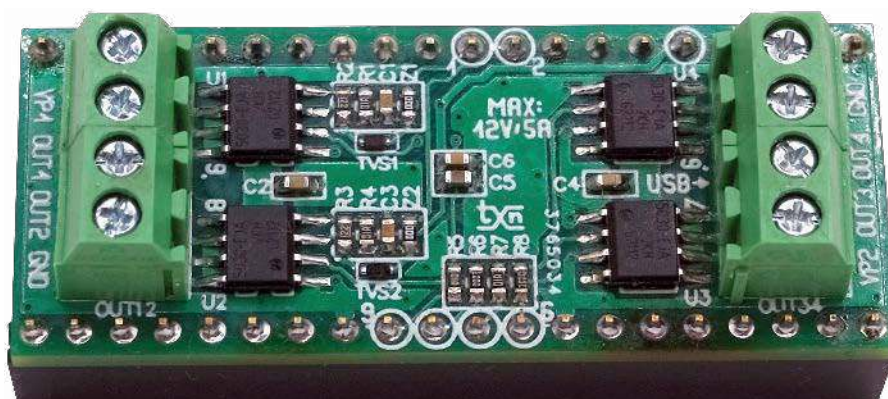
Moduł czterech wyjść HighSide dla RPi Pico

Zaprezentowana płytka rozszerza funkcjonalność Raspberry Pi Pico o cztery wyjścia typu HighSide 12 V/5 A, w tym dwa z możliwością pomiaru prądu obciążenia, co może być przydatne w domowej automatyce czy robotyce.

W module zastosowano specjalizowane klucze półprzewodnikowe HighSide typu BTS5030-1EAJ jako elementy wykonawcze. Ich budowę wewnętrzną pokazuje **rysunek 1**. Oprócz części wykonawczej z zabezpieczeniami mają układ pomiaru prądu wyjściowego przydatny do kontroli stanu obciążenia.

Budowa i działanie

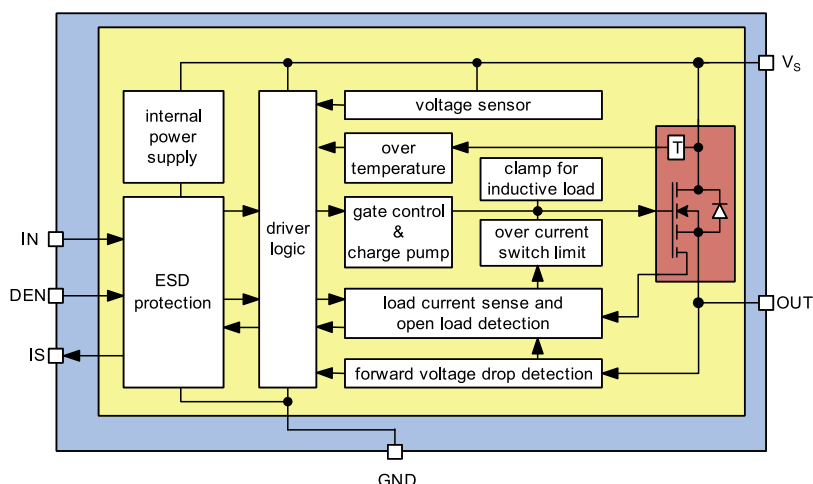
Schemat modułu został pokazany na **rysunku 2**. Klucze sterowane są sygnałami



O1...O4. Zasilanie dołączone na wyprowadzenie VP1 lub VP2 jest kierowane do obciążenia poprzez złącza OUT12 lub OUT34. Zasilanie części mocy zostało rozdzielone na napięcia VP1 i VP2 dla każdej pary wyjść OUT12 i OUT34, aby zwiększyć elastyczność

układu i dostosować do różnych nietypowych aplikacji. Napięcie VPx powinno mieścić się w zakresie 6...15 V, transilne DZ1 i DZ2 zabezpieczają klucze przed skutkami przepięć.

Dwa pierwsze kanały mają możliwość pomiaru prądu obciążenia. Funkcja ta jest aktywowana poprzez stan wysoki na wejściu DEN. Prąd pomiarowy proporcjonalny do prądu wyjściowego jest uzyskiwany na wyprowadzeniu IS. Rezystory R2 i R3 służą do konwersji prądu monitorującego prąd obciążenia. Współczynnik konwersji wynosi $kilim = I_{obc}/I_{is}$ i jest równy ok. 2150. Przy rezystancji 1,2 kΩ daje to napięcie ok. 2,9 V dla prądu 5 A. Ze względu na tolerancję współczynnika kilim i wpływ warunków zewnętrznych pomiar ma dokładność ok. 10% w zakresie 20...100% prądu obciążenia, co jest wystarczające do kontroli poprawności działania obciążenia. Dokładność może zostać zwiększona poprzez odpowiedni dobór zakresu pomiarowego i po kalibracji zgodnie z wytycznymi w karcie katalogowej.



Rysunek 1. Budowa wewnętrzna BTS5030 (za notą Infineon)

Wykaz elementów: (kupuj na stronie sklep.avt.pl lub osobiście: Warszawa, ul. Leszczyńska 11, tel. +48222578451, e-mail: handlowy@avt.pl)

Rezystory: (SMD0603, 1%)
R1, R4, R5, R6, R7, R8: 100 Ω
R2, R3: 1,2 kΩ

Kondensatory: (SMD0603)
C1, C2, C3, C4: 0,1 μF/50 V ceramiczny
C5: 0,1 μF/10 V ceramiczny

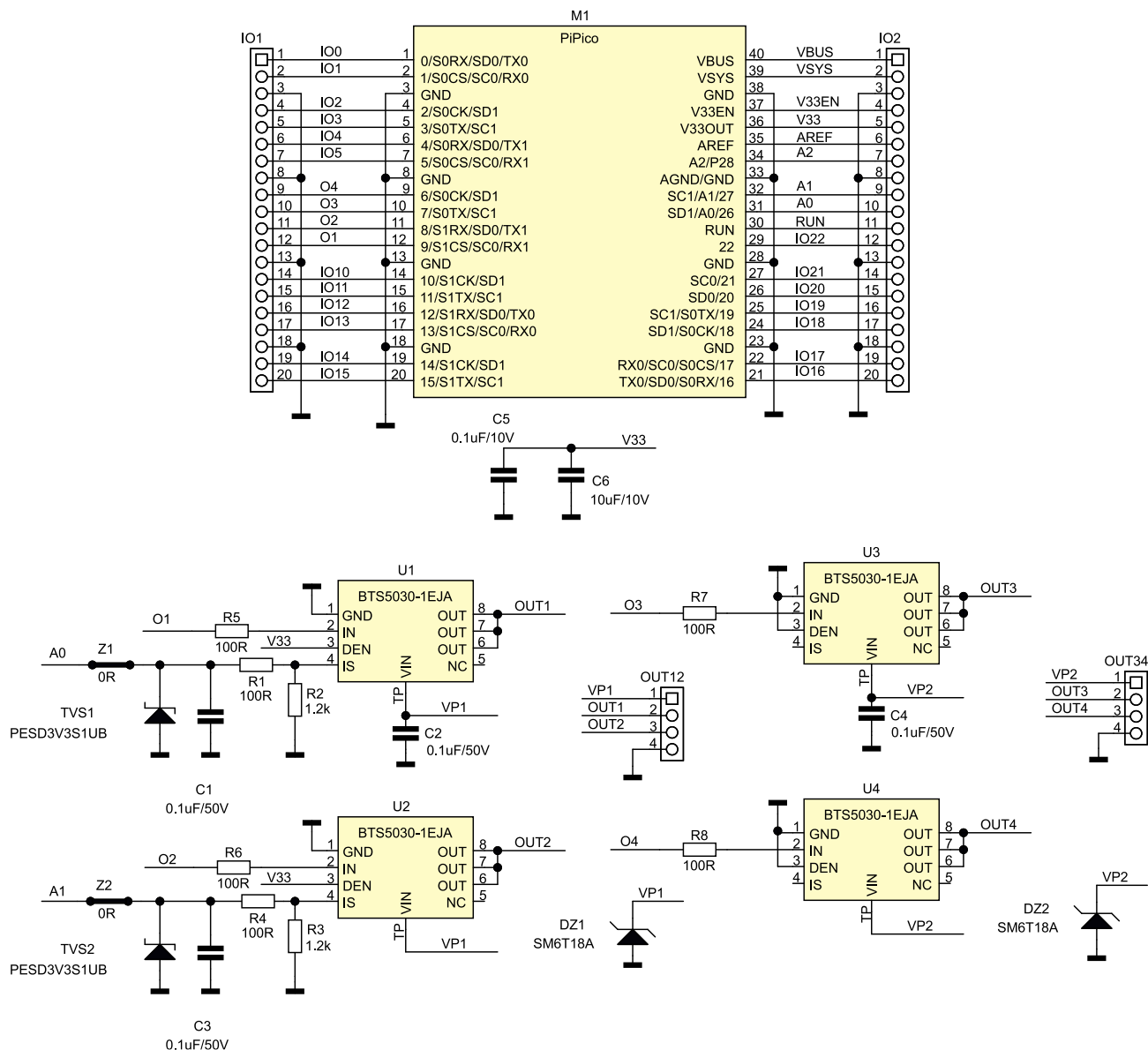
C6: 10 μF/10 V ceramiczny

Półprzewodniki:

DZ1, DZ2: transil jednokierunkowy SM6T18A (SMB_D)
TVS1, TVS2: dioda zabezpieczająca PESD3,3S1UBV (SOD523)
U1, U2, U3, U4: BTS5030-1EAJ (PG-DSO-8)

Pozostałe:

IO1, IO2: gniazdo SIP20 żeńskie
OUT12, OUT34: złącze śrubowe DG 3,5 mm 4 piny (DG381-3.5-4)
Z1, Z2: 0R – opcjonalna zwora SMD (SMD0603)



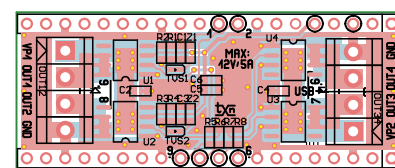
Rysunek 2. Schemat modułu

Montaż i uruchomienie

Układ zmontowany jest na niewielkiej dwustronnej płytce drukowanej, której schemat został pokazany na **rysunku 3**. Montaż nie wymaga opisu, zmontowaną płytkę pokazano na fotografii tytułowej. Dla sprawdzenia

działania można uruchomić prosty skrypt testowy *HiSideRelay.py*, którego kod pokazano na **listingu 1**.

Adam Tatuś, EP



Rysunek 3. Schemat płytki PCB

Listing 1. Skrypt testowy HiSideRelay.py

```
from machine import Pin
from utime import sleep
import time

A0 = machine.ADC(26)
A1 = machine.ADC(27)

O1 = Pin(9, Pin.OUT)
O2 = Pin(8, Pin.OUT)
O3 = Pin(7, Pin.OUT)
O4 = Pin(6, Pin.OUT)

print('01..4=OFF')
O1.value(0)
O2.value(0)
O3.value(0)
O4.value(0)
sleep(1)

print('01=ON')
O1.value(1)
sleep(1)
print('01=OFF')
O1.value(0)
sleep(0.2)

print('02=ON')
O2.value(1)
sleep(1)
print('02=OFF')
O2.value(0)
sleep(0.2)

print('03=ON')
O3.value(1)
sleep(1)
print('03=OFF')
O3.value(0)
sleep(0.2)

print('04=ON')
O4.value(1)
sleep(1)
print('04=OFF')
O4.value(0)
sleep(0.2)

print('01..4=ON')
O1.value(1)
O2.value(1)
O3.value(1)
O4.value(1)

sleep(10)
print('01..4=OFF')
O1.value(0)
O2.value(0)
O3.value(0)
O4.value(0)
sleep(0.2)

Ic0 = A0.read_u16()
print('Prąd 01=OFF ', Ic0)
O1.value(1)
sleep(0.2)
Ic0 = A0.read_u16()
print('Prąd 01=ON ', Ic0)
O1.value(0)
print('01=OFF')

Ic1 = A1.read_u16()
print('Prąd 02=OFF ', Ic1)
O2.value(1)
sleep(0.2)
Ic1 = A1.read_u16()
print('Prąd 02=ON ', Ic1)
O2.value(0)
print('02=OFF')
```

**Podstawowe parametry:**

- pokazywanie jednej z pięciu zapamiętanych ikon emocji: radość, zaskoczenie, obojętność, zmieszanie, smutek,
- przyciski monostabilne przypisane do każdej ikony,
- możliwość wylosowania aktualnie wyświetlanej ikony,
- zapamiętywanie aktualnie wskazywanej ikony na wypadek zaniku zasilania,
- wyświetlacz w postaci dużej matrycy LED 8×8,
- pobór prądu około 70 mA,
- zasilanie napięciem stałym 5 V poprzez gniazdo USB typu B lub listwę zaciskową.

*** Uwaga!** Elektroniczne zestawy do samodzielnego montażu. Wymagana umiejętność lutowania! Podstawową wersją zestawu jest wersja [B] nazywana potocznie KIT-em (z ang. zestaw). Zestaw w wersji [B] zawiera elementy elektroniczne (w tym [UK] – jeśli występuje w projekcie), które należy samodzielnie wylutować w dołączoną płytkę drukowaną (PCB). Wykaz elementów znajduje się w dokumentacji, która jest podlinkowana w opisie kitu. Mając na uwadze różne potrzeby naszych klientów, oferujemy dodatkowe wersje:

- wersja [C] – zmontowany, uruchomiony i przetestowany zestaw [B] (elementy wylutowane w płytkę PCB),
 - wersja [A] – płytkę drukowaną bez elementów i dokumentacji.
- Kity, w których występuje układ scalony wymagający zaprogramowania, mają następujące dodatkowe wersje:
- wersja [A+] – płytkę drukowaną [A] + zaprogramowany układ [UK] i dokumentacja,
 - wersja [UK] – zaprogramowany układ.

Dodatkowe materiały do pobrania ze strony www.ulubionykiosk.pl/media

AVTS670 Pulsujące serce LED (EP 2/2019)
AVT1608 LEDowe serduszko (EP 2/2011)

Nie każdy zestaw AVT występuje we wszystkich wersjach! Każda wersja ma załączony ten sam plik PDF! Podczas składania zamówienia upewnij się, którą wersję zamawiasz! <http://sklep.avt.pl>

W przypadku braku dostępności na stronie sklepu osoby zainteresowane zakupem płytek drukowanych (PCB) prosimy o kontakt via e-mail: kity@avt.pl.

W ofercie AVT*

AVT6014

Wskaźnik stanu emocjonalnego

Jesteś w dobrym nastroju i chcesz to pokazać wszystkim dookoła? A może ogarnia Cię zobojętnienie bądź smutek i lepiej teraz Cię nie zaczepiać? Wydarzyło się coś zaskakującego? Twoje uczucia są... mieszane? Wyraż to! Oczywiście – najlepiej za pomocą elektroniki.

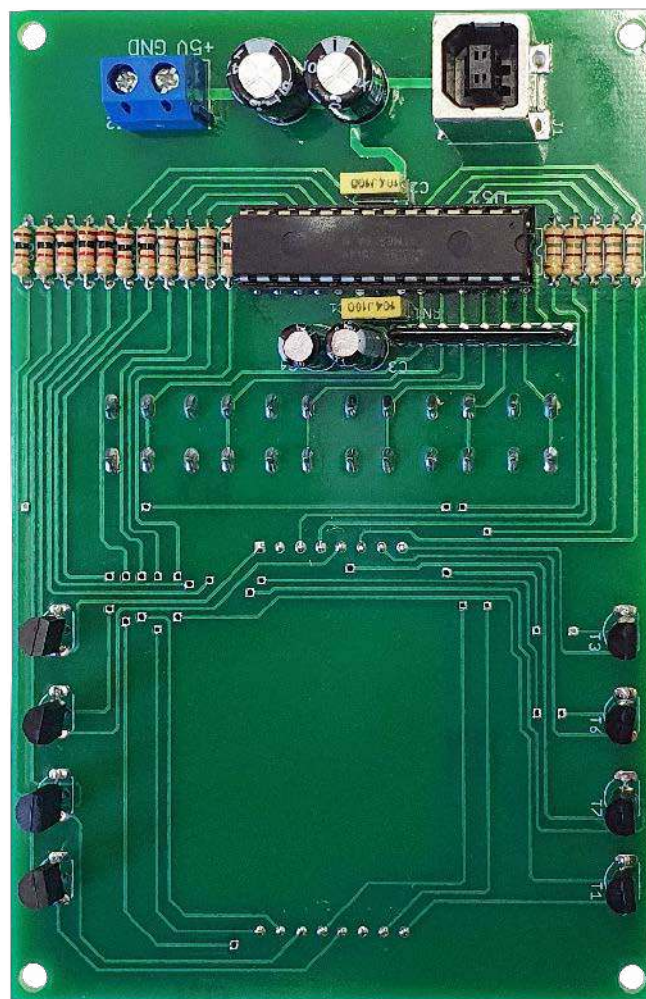
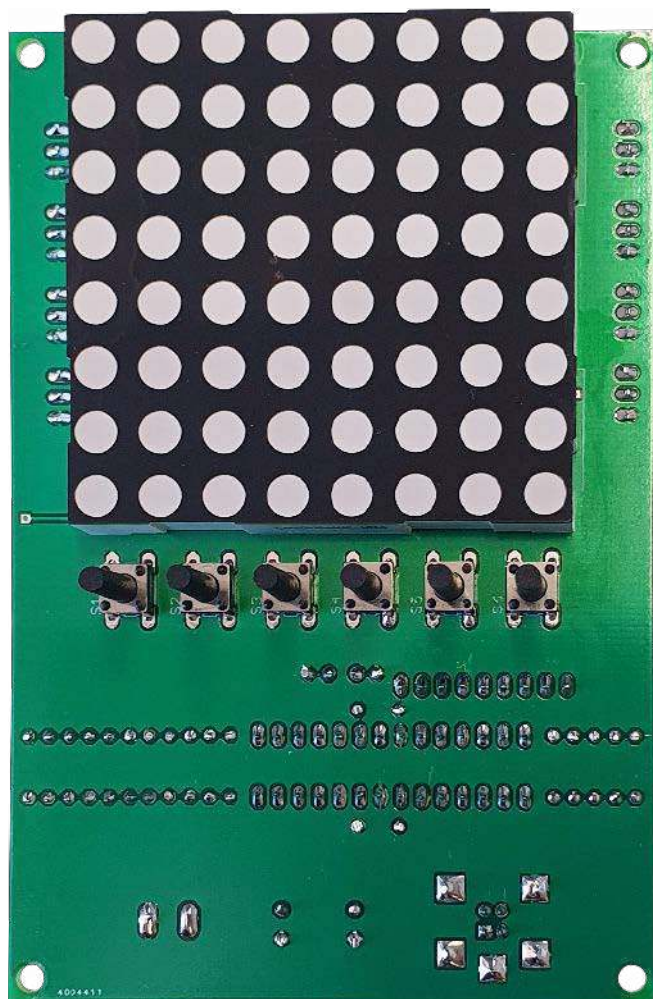
Zaprezentowany układ to typowy gadżet – pokazuje na swoim wyświetlaczu „emotkę” złożoną ze świecących punktów. Można w ten sposób przekazać otoczeniu, w jakim stanie psychicznym aktualnie się znajdujemy. Jeżeli ktoś sam nie wie, co mu w głowie siedzi, może sobie swój stan emocjonalny... wylosować! Została do tego zaimplementowana odpowiednia „maszyna losująca”.

Wyświetlany rysunek jest duży i czytelny, więc widać go nawet z daleka. Może stanowić

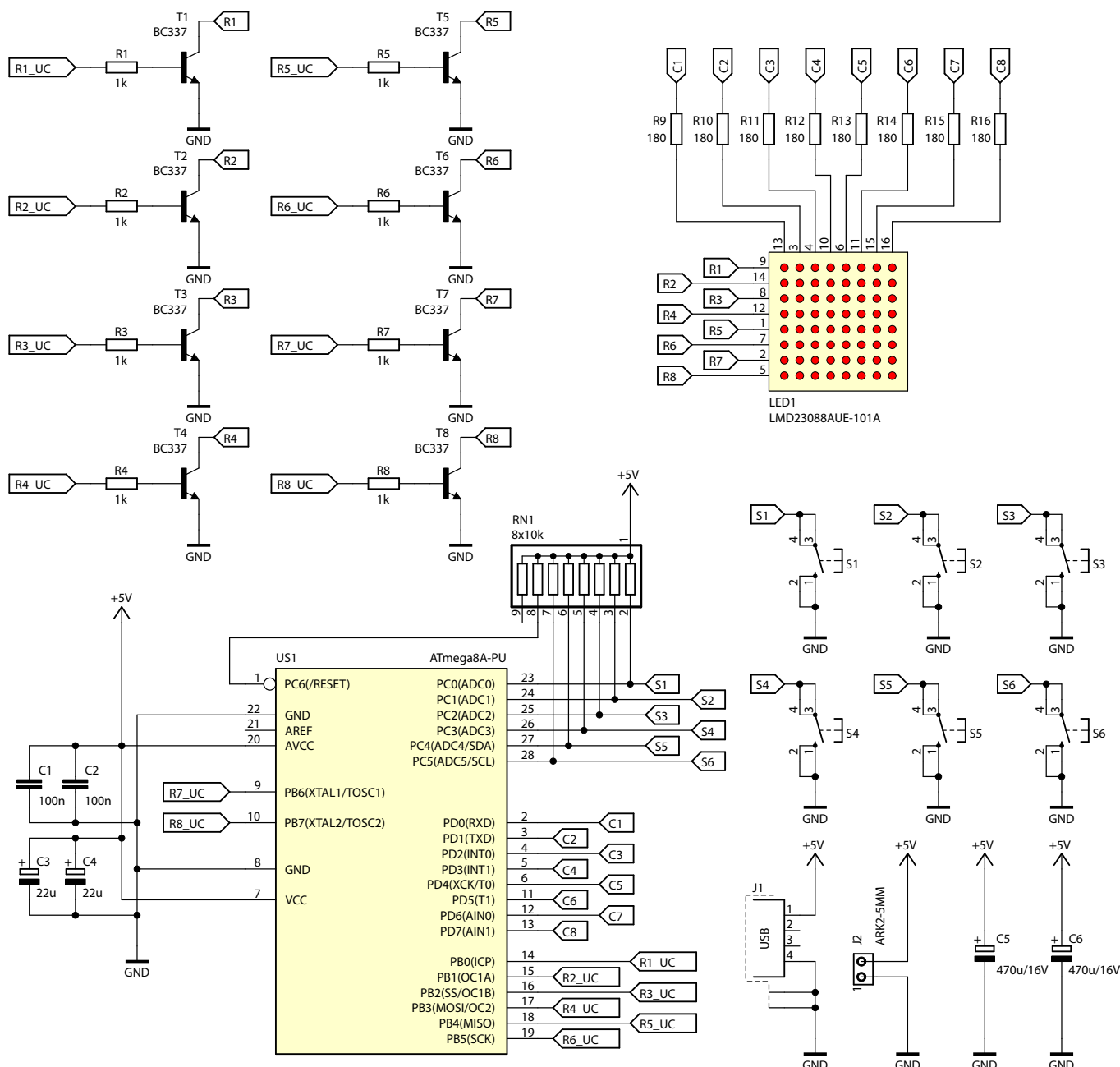
dla otoczenia (np. współpracowników) jasną i czytelną informację na temat tego, czego mogą się po danej osobie w danej chwili spodziewać. Proste? Ale niecodzienne!

Budowa i działanie

Schemat ideowy omawianego układu znajduje się na **rysunku 1**. Głównym podzespołem, zawiadującym jego pracą, jest mikrokontroler typu ATmega8A-PU z 8-bitowym rdzeniem AVR. Ma wystarczającą liczbę konfigurowalnych



Fotografia 1. Wygląd zmontowanej płytki – strona BOTTOM



Rysunek 1. Schemat ideowy wskaźnika stanu emocjonalnego

wyprowadzeń, więc nie zachodzi potrzeba stosowania dodatkowych układów pośredniczących. Dwa jego porty sterują matrycą LED, trzeci odpowiada za odczytywanie stanu styków przycisków monostabilnych. Jego rdzeń jest taktowany wbudowanym generatorem RC o częstotliwości oscylacji 8 MHz – mikrokontroler nie realizuje zadań krytycznych czasowo, zatem taki wzorec częstotliwości jest wystarczający.

Użyta matryca LMD23088AUE-101A ma 64 punkty świejące w kolorze czerwonym,

ułożone w kwadrat 8×8. Anody diod są ułożone w kolumnach, natomiast katody w wierszach. Każdy wiersz jest załączany poprzez nasycenie jednego z tranzystorów NPN T1...T8, natomiast prąd zasilający wszystkie kolumny wypływa bezpośrednio z wyprowadzeń mikrokontrolera i jest ograniczany przez rezystory o wartości 180 Ω. Zapewnia to dostatecznie wysoką jasność świecenia bez ryzyka uszkodzenia układu US1. Jest on odświeżany z częstotliwości 125 Hz: kolejne wiersze są załączane co 1 ms. Wiersze są wybierane

międzyliniowo (1. wiersz – 3. wiersz – 5. wiersz – 7. wiersz – 2. wiersz itd.), a nie kolejno liniowo dla zmniejszenia efektu migotania wyświetlacza.

Użytkownik ma do dyspozycji sześć przycisków monostabilnych S1...S6, którymi zadaje ikonkę do wyświetlania bądź losuje jedną z pięciu zapamiętanych. Rezystory podciągające, wbudowane w mikrokontroler, są wspomagane przez rezystory zawarte w drabince rezystorowej RN1, z których każdy ma rezystancję 10 kΩ. Zmniejsza to wrażliwość

Wykaz elementów: (kupuj na stronie sklep.avt.pl lub osobiście: Warszawa, ul. Leszczyńska 11, tel. +48222578451, e-mail: handlowy@avt.pl)

Rezystory: (THT o mocy 0,25 W):

R1...R8: 1 kΩ
R9...R16: 180 Ω
RN1: 8 × 10 kΩ SIL9

Kondensatory:

C1, C2: 100 nF raster 5 mm MKT

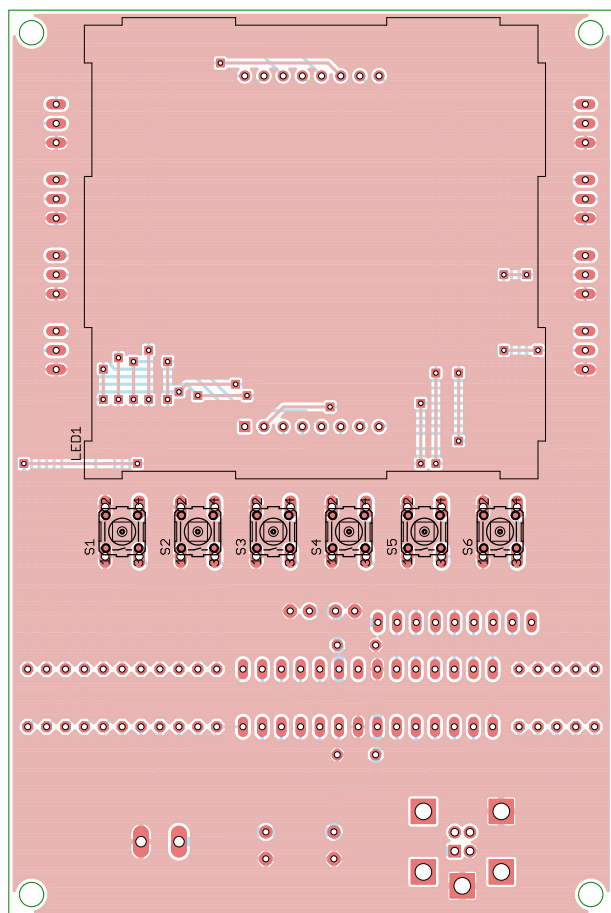
C3, C4: 22 μF 25 V raster 2,5 mm
C5, C6: 470 μF 16 V raster 3,5 mm

Półprzewodniki:

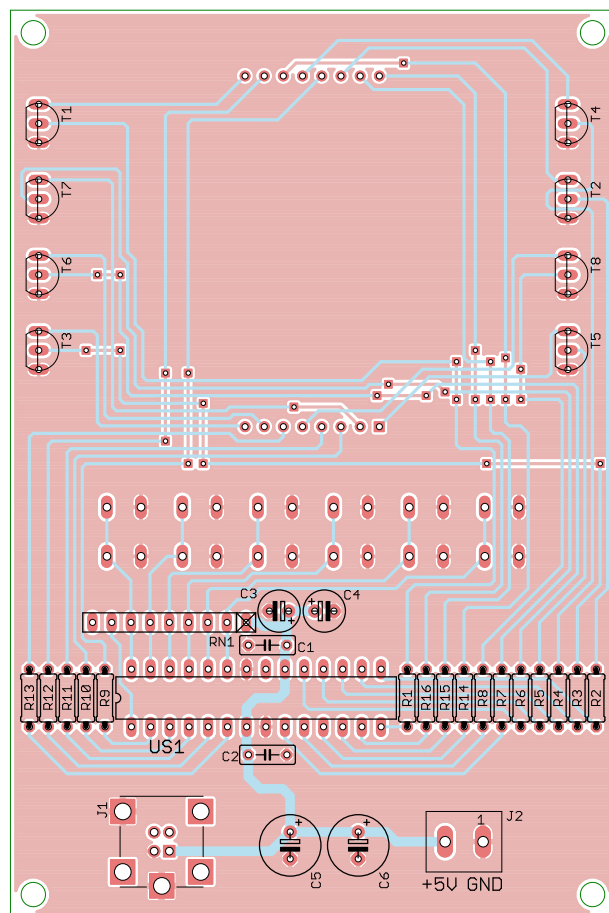
LED1: LMD23088AUE-101A (opis w tekście)
T1...T8: BC337 TO92
US1: ATmega8A-PU DIP28

Pozostałe:

J1: UBBS-4R-D14-4D
J2: ARK2/500
S1...S6: microswich 6×6 13,5 mm
Jedna podstawka DIP28 wąska



Rysunek 2. Schemat płytki PCB



Rysunek 3. Ustawienie bitów zabezpieczających

układu na zakłócenia elektromagnetyczne. Linia zerująca mikrokontrolera również została podciągnięta do dodatniego potencjału zasilającego w tym samym celu.

Na płytce znalazło się gniazdo USB typu B, którym można zasilac układ, jak również listwa zaciskowa. Pochodzące z zewnątrz napięcie stałe o wartości 5 V jest filtrowane przez łącznie sześć kondensatorów o zróżnicowanej pojemności, dla lepszego odsprężania zasilania dla mikrokontrolera.

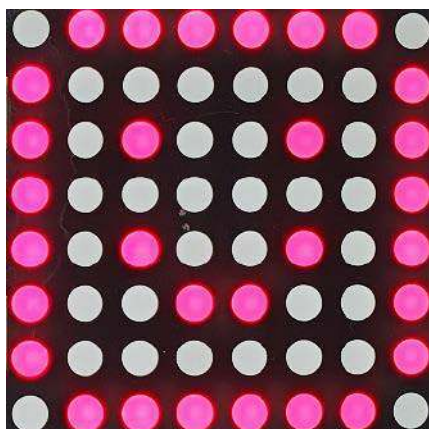
Montaż i uruchomienie

Układ został zmontowany na dwustronnej płytce drukowanej o wymiarach 120×80 mm. Jej wzór ścieżek oraz schemat montażowy zostały pokazane na **rysunku 2**.

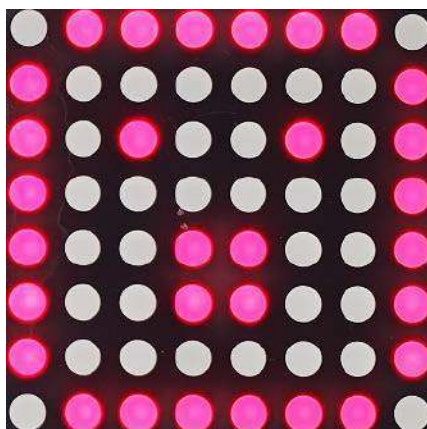
W odległości 3 mm od krawędzi płytki znalazły się cztery otwory montażowe, każdy o średnicy 3,2 mm.

Montaż proponuję rozpocząć od elementów o najmniejszej wysokości obudowy, czyli rezystorów. Pod układ US1 proponuję zastosować podstawkę, aby ułatwić jego programowanie oraz wymianę w razie uszkodzenia. W układzie prototypowym wyświetlacz LED1 oraz przyciski S1...S6 znalazły się na wierzchniej stronie płytki (TOP), zaś cała reszta elementów na stronie spodniej (BOTTOM). Zmontowany układ można zobaczyć na fotografii tytułowej oraz **fotografii 1**. Nic nie stoi na przeszkodzie, by przyciski wlotować od drugiej strony laminatu.

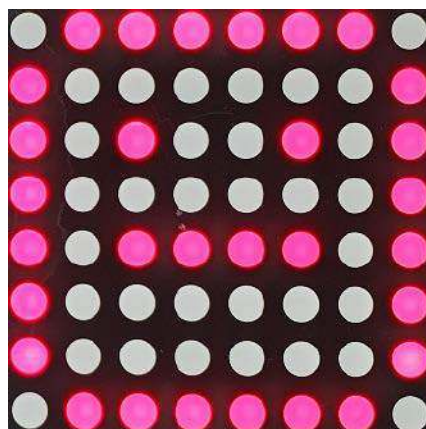
Na etapie uruchamiania konieczne jest zaprogramowanie pamięci Flash mikrokontrolera dostarczoną wsadą oraz zmiana jego bitów zabezpieczających. Oto ich nowe wartości:



Fotografia 2. Ekran nr 1 – buźka uśmiechnięta



Fotografia 3. Ekran nr 2 – buźka zaskoczona



Fotografia 4. Ekran nr 3 – buźka neutralna

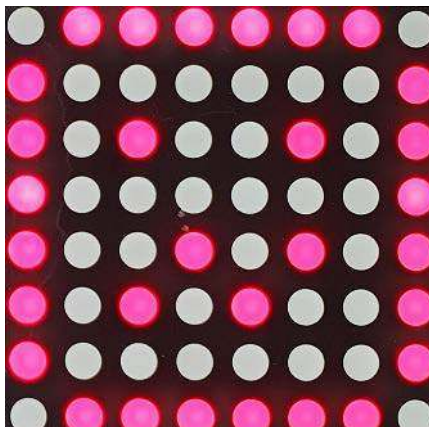
Low Fuse = 0x24

High Fuse = 0xD9

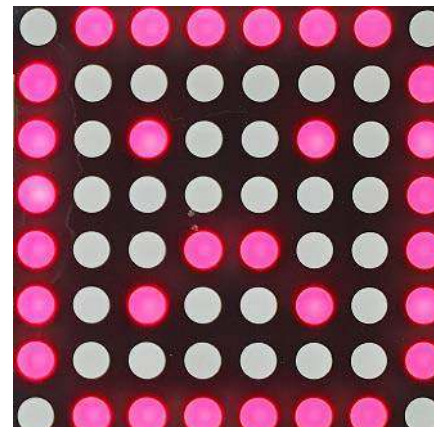
Szczegóły są widoczne na **rysunku 2**, który pokazuje widok okna konfiguracji tych bitów z programu BitBurner. W ten sposób zostanie uruchomiony wewnętrzny generator RC o częstotliwości oscylacji 8 MHz oraz Brown-Out Detector, który wprowadzi mikrokontroler w stan zerowania, jeżeli jego napięcie zasilające spadnie poniżej 4 V. To znacznie zmniejsza ryzyko zawieszenia się mikrokontrolera podczas uruchamiania.

Poprawnie zaprogramowany układ jest gotowy do działania po podłączeniu zasilania do zacisków złącza J1 lub J2. Powinno to być napięcie stałe, dobrze filtrowane, najlepiej stabilizowane. Może pochodzić, na przykład, z ładowarki impulsowej USB. Nominalna wartość tego napięcia powinna wynosić 5 V, lecz dopuszczalne są pewne odstępstwa. Maksymalna wartość to 5,5 V i wynika z ograniczeń nałożonych przez notę katalogową producenta mikrokontrolera, zaś dolną granicę można oszacować na 4,5 V – tak, by zabezpieczenie BOD było dalekie od zadziałania.

Po włączeniu zasilania pokaże się ostatnio zapamiętana „buźka” – jeżeli mikrokontroler



Fotografia 5. Ekran nr 4 – buźka zmieszana



Fotografia 6. Ekran nr 5 – buźka zasmucona

nie był wcześniej włączany, będzie ona wesoła. Każdemu z przycisków S1...S5 odpowiada dokładnie jedna ikona emocji:

- wesoła jest na **fotografii 2**, załączana wciśnięciem przycisku S1,
- zaskoczona widnieje na **fotografii 3**, można ją załączyć, wciskając S2,
- obojętna (**fotografia 4**) pokaże się po zwarciu styków S3,
- zmieszana, jaką można zobaczyć na **fotografii 5**, ukaże się po wciśnięciu S4,
- smutna – **fotografia 6** – wyświetlona po naciśnięciu S5.

Ostatni przycisk, S6, służy do wylosowania jednej z tych pięciu ikon. W czasie jego trzymania ekran jest wygaszony, a układ przeskakuje po wszystkich „emotkach” po kolei, z bardzo wysoką częstotliwością. Ponieważ czas trzymania styków palcem w stanie zwarcia jest losowy, to i pokazaną buźkę można uznać za wylosowaną. Każda z emotek jest zapamiętywana w nieulotnej pamięci EEPROM zaraz po jej wyświetleniu, po czym zostaje przywrócona po załączeniu zasilania.

Michał Kurzela, EP

REKLAMA

www.ep.com.pl/EPwtoku

Czytaj artykuły
zanim zostaną
wydane
w formie
papierowej





Bezpieczne złącza do magazynów energii

Magazynowanie energii to niezwykle istotny temat w branży energoelektroniki w ostatnim czasie. Wiele firm ma już stosowne aplikacje w swoim portfolio, a pozostałe z uwagą przyglądają się zapotrzebowaniu na rynku, aby w porę zareagować i dołączyć do wyścigu w tej obiecującej branży.

Magazynowanie energii przez najbliższe lata będzie tematem wszechobecnym nie tylko bezpośrednio w branży energetycznej i elektronicznej, ale także w mediach, a nawet w rozmowach polityków – w końcu bezpieczeństwo energetyczne kraju to ważny aspekt. Również w rozważaniach prywatnych użytkowników przydomowych instalacji fotowoltaicznych magazyny energii stają się coraz bardziej istotne – za kilka lat zakończy się okres rozliczeń na zasadzie net meteringu i trzeba będzie znaleźć korzystny sposób na zużywanie wyprodukowanej na dachu energii. Inwestorzy chcąc uniezależnić się choć w części od znacznego wzrostu cen energii, będą poszukiwać możliwości zgromadzenia jej zapasu bezpośrednio u siebie.

W artykule skupimy się jednak na bezpiecznym użytkowaniu magazynu energii. Poza ceną kompletnego systemu, od tego będzie zależała mniejsza lub większa popularność rozwiązania danego producenta.

Bezpieczne złącza

Może się wydawać, że zrealizowanie połączenia elektrycznego to banalna sprawa. Jest wtyczka z kablem, jest gniazdo – łączymy oba i gotowe. Spróbujmy jednak rozłożyć to na czynniki pierwsze i pokazać potencjalne punkty krytyczne. Skupmy się w pierwszej kolejności na początkowym montażu przed uruchomieniem.

Magazyn energii zazwyczaj składa się z kilku modułów/pakietów bateryjnych, które łączymy ze sobą (**fotografia 1**). Intuicyjnie oczywiście widać, która wtyczka powinna być połączona z którym gniazdem (kodowanie kolorami). Jednak czy to wystarczy? Nie zawsze. Czasami instalacja może odbywać się w złych warunkach oświetleniowych, nigdy nie wiemy, na ile sprawny jest wzrok oraz jak wygląda kwestia umiejętności i staranności instalatora, który podjął się tego zlecenia. Już na tym etapie istnieje (co prawda niewielkie, ale zawsze możliwe) ryzyko pomyłki z gubnej w skutkach (np. uszkodzenie modułów).

Zdecydowanie większe bezpieczeństwo oferują złącza z kodowaniem mechanicznym, które uniemożliwiają wykonanie niewłaściwego połączenia. Nawet jeśli dobierzemy elementy z kodowaniem mechanicznym, to także możemy napotkać pewne komplikacje. Okazuje się, że nie każdy system dostępny na rynku pozwala na dowolne pozycjonowanie wtyku względem gniazda podczas łączenia! Niektóre rozwiązania pozwalają na połączenie wyłącznie z przewodem kątowym wyprowadzonym w dół. Oczywiście po podłączeniu można obrócić



Fotografia 1. Przykład łączenia szeregowego modułowych pakietów bateryjnych w magazynie energii

kabel tak, aby był skierowany w stronę kolejnego gniazda, ale tam znowu napotykamy ten sam problem – aby wykonać połączenie, przewód musi być skierowany w dół.

Problem wydaje się błahy – wystarczy wykonać dłuższe kable łączeniowe z zapasem na odpowiednie wygięcie całości. Jednak istnieje co najmniej kilka wad takiego rozwiązania. Po pierwsze – dłuższy kabel to więcej miedzi (często wymagane są grube przekroje z racji płynących dużych prądów). Więcej miedzi to większy koszt. Przemnożony przez tysiące kabli, jakie mamy wyprodukować i sprzedać, skutecznie wpławi w zdenerwowanie dział finansowy firmy lub co najmniej od razu wskaże możliwe miejsce wygenerowania oszczędności. Po drugie – grube kable (np. przekroju 90 mm² czy 120 mm²) wcale nie tak łatwo pozwalają się wyginać – trzeba w to włożyć sporo siły. I w końcu po trzecie – skoro musimy w coś wkładać dużo siły, to pośrednio narażamy niektóre miejsca konstrukcji na zbędne (a czasem nieprzewidziane na etapie projektowania) naprężenia mogące prowadzić do powstania uszkodzeń. A jak wiemy z doświadczeń, nawet pozornie niewielkie pęknięcia z czasem mogą wygenerować spore problemy. Nam przecież zależy na niezawodności działania systemu przez jak najdłuższy czas.

Dodatkowe ryglowanie złączy

Kolejna do rozważenia kwestia to bezpieczne zaryglowanie miejsca łączenia. Jak powszechnie wiadomo, każde „luźne połączenie” kiedyś stanie się źródłem mniej lub bardziej groźnych komplikacji. W zależności od typu danego obwodu, prądów, z jakimi mamy do czynienia i napięcia roboczego, może powstawać miejscowe przegrzanie, iskry i wypalanie styków.



Chcesz tworzyć magazyny energii? Postaw na solidne rozwiązania!

Złącza wysokoprądowe do pakietów bateryjnych to kluczowe elementy, od których zależy bezpieczeństwo użytkowania i niezawodność twojego systemu.



Wybierz mądrze i zaufaj ekspertom w branży połączeń!

Odwiedź naszą stronę internetową i dowiedz się więcej!

<https://phoe.co/ess>

lub zeskanuj kod QR swoim telefonem



REKLAMA



Producenci tak projektują metalowe styki, aby po ich połączeniu uzyskać jak najniższą rezystancję przy możliwie małych siłach dociskowych (zapewnienie łatwości obsługi i trwałości powłok galwanicznych). Dodatkowe ryglowanie na obudowie złącza przejmując te funkcje. Z tego powodu warto postawić na rozwiązania zaprojektowane w zgodzie z ideą Poka Yoke, czyli w tym wypadku „połącz w łatwy sposób, bez zbędnego komplikowania operacji, wymagającego dodatkowych instrukcji” (fotografia 2).

W najlepszych rozwiązaniach wystarczający jest zaledwie jeden ruch (wcisnięcie wtyku na gniazdo z automatycznym ryglowaniem). Do odłączenia natomiast wystarczy przesunąć umieszczony w konstrukcji wtyczki przycisk i zdjąć końcówkę kabla z gniazda. Brak dodatkowych elementów to niezaprzeczalna zaleta tego rozwiązania. Spotykane wykonania niektórych producentów mają np. dodatkową dźwignię ryglującą, która wymaga zamknięcia. Niestety pozostawia to miejsce na błąd instalatora, który w pośpiechu lub z niestaraności może pominąć ten etap. Konsekwencje? Ryzyko poluzowania połączenia wskutek wibracji lub niespodziewanych/przypadkowych szarpnięć kablami, prowadzące do wcześniej opisanych skutków.

Istotne jest również zabezpieczenie elementów będących pod napięciem przed przypadkowym dotknięciem palcem czy zwarcie np. o metalową obudowę urządzenia. Tu należy wybierać wyłącznie takie konstrukcje, które mają stosowne zabezpieczenie (fotografia 3).

Montaż wtyku i gniazda

Z pozoru łatwa czynność, taka jak montaż wtyku na przewodzie, może czasami okazać się problematyczna. Zazwyczaj w aplikacjach tego typu (przy stosunkowo dużych przekrojach i prądach) jako najbardziej niezawodne i trwałe stosuje się połączenia zaciskane. Dobrze, jeśli sam proces przygotowania kabla oraz zaprasowywania nie wymaga zbyt wielu operacji, jak i same złącza nie składają się ze zbyt wielu elementów. Posługując się przykładem – wszelkie dodatkowe metalowe tulejki, które musimy zakładać pośrednio na odizolowany przewód przed umieszczeniem go w komorze zaciskowej złącza, to elementy, które mogą zostać przeoczone (lub zagubione). A musimy sobie zdawać sprawę, że producenci,



Fotografia 2. Warto zwrócić uwagę na to, czy system kodowania pozwala na zorientowanie przewodu w dowolną/wygodną stronę



Fotografia 3. Odpowiednie wykonanie zarówno gniazda, jak i wtyku gwarantuje odpowiednie ukrycie wewnątrz obudów odizolowanych podzespołów złączy. Dzięki temu zarówno instalator, jak i sam sprzęt nie jest narażony na niebezpieczne sytuacje

projektując swoje komponenty, przeznaczą je do montażu w konkretny sposób, który gwarantuje niezawodność. Lepiej więc wybierać to, co jest łatwiejsze w montażu i nie pozostawia możliwości popełnienia błędów.

Należy również mieć na uwadze sam montaż gniazda na obudowie urządzenia. Czy bezpieczniejsze (czytaj: bardziej odporne na ewentualne zniszczenia) jest przykręcenie kwadratowej flanszy gniazda w każdym narożniku, czy sytuacja, gdzie cała izolacja elementu trzymana jest na jednym sworzniu? W przypadku uszkodzenia mechanicznego gniazda możemy ryzykować odpadnięcie całej ramki złącza i całkowite odsłonięcie elementu pod napięciem, versus ewentualne odprysnięcie fragmentu tworzywa.

Podsumowanie

Projektanci urządzeń mają wiele aspektów do przeanalizowania i uwzględnienia przy doborze odpowiednich podzespołów. Z tego powodu warto polegać na firmach, które mają wieloletnie doświadczenie w danej dziedzinie (np. systemów połączeń). Są one gwarantem nie tylko właściwej jakości, ale także przemyślanej konstrukcji nawet nowo powstających rozwiązań, wychodzącej naprzeciw oczekiwaniom użytkowników (często nawet nieświadomych konkretnych zagrożeń). Jak to w życiu zwykle bywa, „diabeł tkwi w szczegółach” – nie zawsze to, co na pierwszy rzut oka wygląda podobnie, ma tak samo dopracowaną funkcjonalność (mowa tu o tanich kopiach markowych produktów). Niezawodność, bezpieczeństwo instalatorów czy użytkowników – to już tylko bonusy (ale jakże istotne), które mogą przeważać o sukcesie finalnego produktu na rynku.

Jeśli interesuje Cię powyższa tematyka i szukasz pomocy w doborze właściwego rozwiązania – zapraszamy do kontaktu. Napisz, zadzwoń, skorzystaj z naszej strony internetowej (PHOENIX CONTACT, Złącza do zasobników energii).

Piotr Andrzejewski
PHOENIX CONTACT



Maksymalizuje sprawność elektroniki

Kontakt Chemie od ponad 60 lat jest czołowym producentem aerozoli technicznych na potrzeby przemysłu elektronicznego. W tym czasie firma wkroczyła poza sektor elektroniki, aby opracować asortyment wysokiej jakości produktów na potrzeby wszelkiego rodzaju sprzętów elektrycznych i elektronicznych.

Produkt Kontakt 60 opracowany przez firmę Kontakt Chemie od samego początku był przełomową innowacją i odniósł ogromny sukces w Niemczech. Od tego czasu firma wprowadziła na rynek również wiele innych artykułów, które sprawiają, że jej reputacja na polu skuteczności i wydajności stale rośnie. Do tego stopnia, że trudno wyobrazić sobie dzisiaj konstruowanie, konserwację i naprawę sprzętów elektrycznych i elektronicznych bez użycia produktów Kontakt Chemie.

Tlenki nie mają szans

Kontakt 60 to środek do czyszczenia styków rozpuszczający tlenki. Jest szczególnie zalecany na potrzeby regeneracji skorodowanych, mocno zużytych i zabrudzonych styków. Rozpuszcza tlenki na stykach i przywraca niską rezystancję połączenia, co gwarantuje niski spadek napięcia nawet przy najmniejszych naciskach na styk.

Oczyszczone obszary można od razu spłukać przy użyciu preparatu Kontakt WL. Natomiast dodatkową ochronę zapewnia Kontakt 61. Smaruje czyste styki, nie oddziałuje na metale, grafit, materiały na bazie węgla, tworzywa termoplastyczne, żywice termoutwardzalne, izolatory i inne materiały. Jest dielektrykiem, który nie stymuluje prądów upływowch.

Kontakt 60 zapewnia doskonałe rezultaty wszędzie tam, gdzie należy wyczyścić styki elektryczne. Usuwa tlenki z wszelkiego rodzaju styków metalowych w elektronice, motoryzacji i zastosowaniach przemysłowych, takich jak przełączniki, zespoły wtyczkowe, gniazda układów scalonych, bezpieczniki, oprawy lamp, styki ślizgowe.

Użyj trzyetapowego procesu czyszczenia styków Kontakt Chemie, aby uzyskać wiarygodne wyniki:

- Kontakt 60 – rozpuszcza tlenki i regeneruje mocno zużyte styki;
- Kontakt WL – spłukuje rozpuszczone pozostałości korozji, tłuszcz oraz brud i całkowicie odparowuje;
- Kontakt 61 – cienka powłoka chroniąca przed korozją i zużyciem, wydłużająca czas konserwacji i niezawodność.

CRC Industries Europe BV
www.kontaktchemie.com

**KONTAKT
CHEMIE®**

By CRC Industries **CRC**

3-ETAPOWA KONSERWACJA STYKÓW

1 ROZPUŚĆ TLENKI ZA POMOCĄ
KONTAKT 60

OCZYŚĆ/ODTŁUŚĆ ZA POMOCĄ
KONTAKT WL

2

3 ZABEZPIECZ ZA POMOCĄ
KONTAKT 61



REKLAMA



Skazani na lit

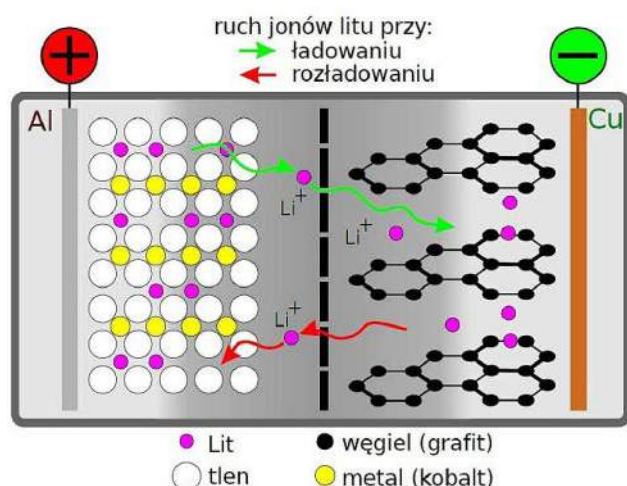
Energię trudno magazynować, zwłaszcza na dłuższy czas. Akumulatory litowe dają taką możliwość i w przeciwieństwie do wielu wcześniejszych rozwiązań są wolne od takich wad jak efekt pamięciowy, samorozładowanie czy niska żywotność. Jednak najważniejszą cechą, która przyczyniła się do tego, że zdominowały rynek, jest wysoka gęstość energii – najwyższa spośród wszystkich akumulatorów dostępnych na rynku.

Niestety technologia akumulatorowa nie osiągnęła takiego samego wzrostu wydajności jak np. technika cyfrowa. Pomimo ciągłych innowacji i ulepszeń akumulatory nie są w stanie sprostać wielu stawianym im wymaganiom. Jednak intensywność badań nie ustaje i stale ogłaszane są nowe osiągnięcia i ulepszenia. Choć baterie litowe są dziś bardziej popularne niż kiedykolwiek wcześniej, to możemy nie zdawać sobie sprawy z tego, że jest ich już kilka typów. Jednak zacznijmy od początku.

Budowa i zasada działania

W akumulatorze litowym podstawą magazynowania energii jest ruch dodatnich jonów litu między anodą i katodą w przewodzącym elektrolicie, co związane jest z przemianami chemicznymi. Dziś dostępnych jest kilka odmian ogniw litowych. Mają różne nazwy ale wszystkie są akumulatorami litowo-jonowymi. Dodatnia elektroda wykonana jest z różnych związków litu, które mogą oddawać i przyjmować jony litu. Natomiast elektroda ujemna (anoda) zbudowana jest najczęściej z jakiejś odmiany węgla, choć niektóre odmiany mają anody z innych materiałów. Elektrody muszą być od siebie odseparowane, żeby nie nastąpiło zwarcie, natomiast między nimi musi być umieszczony elektrolit, który pozwoli na przemieszczanie się dodatnich jonów litu między elektrodami. Uproszczoną konstrukcję ogniwa litowego pokazano na **rysunku 1**.

Przez lata zbadano najróżniejsze związki litu, aby wyszukać takie, które łatwo oddają i przyjmują te jony. Okazało się, że do budowy katody nadają się związki litu, mające specyficzną budowę kryształową. Przeważnie są to tlenki metali, ale nie jednego metalu (litu), tylko



Rysunek 1. Uproszczona budowa i zasada działania akumulatora litowego



Fotografia 1. Miniaturowy akumulator LCO o pojemności 1500 mAh

dwóch, a nawet trzech. W pierwszych praktycznie użytecznych akumulatorach litowych (Sony 1991) katoda zbudowana była z tlenku litu i kobaltu – LiCoO_2 .

Akumulatory litowo-kobaltowe (LiCoO_2 – LCO)

Akumulator składa się z katody z tlenku kobaltu i grafitowej anody węglowej. Katoda ma budowę warstwową i podczas rozładowywania jony litu przemieszczają się od anody do katody. Przepływ odwraca się po naładowaniu.

Ogniwa mają dobrą zdolność magazynowania energii – wysoką energię właściwą na poziomie 150...200 Wh/kg, a specjalne ogniwa zapewniają do 240 Wh/kg. Oznacza to, że mogą dostarczać energię przez długi czas. Dostępne są jako standardowe ogniwa cylindryczne różnego typu – **fotografia 1**. Napięcie maksymalne wynosi zwykle 4,20 V. Niestety mają szereg wad:

- stosunkowo krótka żywotność – zwykle od 500 do 1000 cykli,
- niska stabilność termiczna – co prowadzi do obaw o bezpieczeństwo,
- ograniczona obciążalność (moc właściwa)
- kobalt jest dość drogi.

Ogniwo nie należy ładować ani rozładowywać prądem wyższym niż wartość znamionowa C. Oznacza to, że ogniwo 18650 o pojemności 2400 mAh można ładować i rozładowywać prądem o maksymalnej wartości 2400 mA. Wymuszanie szybkiego ładowania lub przyłożenie obciążenia większego niż C powoduje przegrzanie. Aby zapewnić optymalne szybkie ładowanie, przyjmuje się współczynnik 0,8 C. Obwód zabezpieczający akumulator powinien ograniczać szybkość ładowania i rozładowywania do bezpiecznego poziomu.

Z uwagi na koszt wytwarzania, z początku budowano jedynie akumulatory LiCoO_2 o małych rozmiarach, przeznaczone do kosztownego sprzętu elektronicznego. LCO traci przychylność na rzecz NMC i NCA.

Akumulatory litowo-niklowo-manganowo-kobaltowe (LiMnCoO_2 – NMC)

Lepsze od kobaltowych okazały się kryształy tlenku zawierające atomy litu oraz manganu i kobaltu LiMnCoO_2 (LMC), a jeszcze bardziej niklu, manganu i kobaltu LiNiMnCoO_2 (NMC). Zastosowanie takich materiałów pozwalało optymalizować parametry akumulatorów do rozmaitych zastosowań:

- NMC w ogniwie 18650 przy umiarkowanym obciążeniu ma pojemność około 2800 mAh i może dostarczać prąd od 4 A do 5 A.
- NMC w tym samym ogniwie zoptymalizowanym pod kątem określonej mocy ma pojemność tylko około 2000 mAh, ale zapewnia ciągły prąd rozładowania o natężeniu 20 A.
- Konstrukcja z anodą na bazie krzemu pozwala osiągnąć 4000 mAh i więcej, ale przy zmniejszonej obciążalności i krótszym cyklu życia. Krzem dodany do grafitu ma tę wadę, że anoda rośnie i kurczy się pod wpływem ładowania i rozładowania, co powoduje, że ogniwo jest niestabilne mechanicznie.



Fotografia 2. Akumulator litowy typu NMC



Fotografia 3. W autach Tesla często stosowane ogniwa to NMC i NCA

Sekret NMC tkwi w połączeniu niklu i manganu. Nikiel jest znany z wysokiej energii właściwej, ale słabej stabilności; mangan ma tę zaletę, że tworzy strukturę spinelową, aby osiągnąć niski opór wewnętrzny, ale oferuje niską energię właściwą. Łączenie metali wzmacnia wzajemne mocne strony.

NMC to akumulatory wybierane do elektronarzędzi, rowerów elektrycznych i innych elektrycznych układów napędowych (**fotografia 2**). Kombinacja katod składa się zazwyczaj z jednej trzeciej niklu, jednej trzeciej manganu i jednej trzeciej kobaltu, znanej również jako 1-1-1. Kobalt jest drogi i ma ograniczoną podaż. Producenci akumulatorów zmniejszają zawartość kobaltu, rezygnując z pewnego kompromisu w zakresie wydajności. Udaną kombinacją jest NCM532 z 5 częściami niklu, 3 częściami kobaltu i 2 częściami manganu. Inne kombinacje to NMC622 i NMC811.

NMC mają wyższą stabilność termiczną niż akumulatory LCO, co czyni je ogólnie bezpieczniejszymi. Główną wadą baterii NMC jest to, że mają nieco niższe napięcie niż baterie na bazie kobaltu. Nowe elektrolity i dodatki umożliwiają ładowanie do napięcia 4,4 V/ogniwo i wyższego, aby zwiększyć pojemność.

Akumulatory litowo-niklowo-kobaltowo-aluminiowe (LiNiCoAlO_2 – NCA)

Oferują wysoką energię właściwą, przyzwolą moc właściwą i długi cykl życia. Oznacza to, że mogą dostarczać stosunkowo dużą ilość prądu przez dłuższy czas. Te cechy sprawiły, że baterie NCA są popularne na rynku pojazdów elektrycznych. W autach Tesla często stosowane ogniwa to NMC i NCA – (**fotografia 3**).

Badania sugerują, że akumulatory NCA mogą wytrzymać w odpowiednich warunkach ponad 40 lat. Zawartość aluminium nie tylko pomaga zapobiegać zniszczeniu, ale także pozwala, aby katoda zawierała aż 84% niklu. Dzięki temu ich produkcja jest znacznie tańsza w porównaniu z innymi akumulatorami bogatymi w nikiel.

Akumulatory litowo-manganowe (LiMn_2O_4 – LMO)

Ogniwo litowo-jonowe z tlenkiem litowo-manganowym jako materiałem katody powstało już w 1996 r. Architektura tworzy trójwymiarową strukturę spinelową, która poprawia przepływ jonów na elektrodzie, co skutkuje niższym oporem wewnętrznym i lepszą wydajnością prądową. Kolejną zaletą spinelu jest wysoka stabilność termiczna i zwiększone bezpieczeństwo, ale żywotność wypada gorzej od poprzednich ogniw – ok. 300...700 cykli.

Niska rezystancja wewnętrzna LMO umożliwia szybkie ładowanie i rozładowywanie wysokim prądem. Pakiet 18650 Li-mangan może być rozładowywany prądem 20...30 A przy umiarkowanym nagrzewaniu się. Możliwe jest także zastosowanie jednosekundowych impulsów obciążających do 50 A. Ciągłe wysokie obciążenie tym prądem spowodowałoby



Fotografia 4. Akumulatory LMO stosowane w autach Nissan Leaf

gromadzenie się ciepła, a temperatura ogniwa nie może przekroczyć 80°C. LMO stosowane są w elektronarzędziach, instrumentach medycznych, a także pojazdach hybrydowych i elektrycznych.

Ogniwo LMO ma pojemność mniej więcej o jedną trzecią niższą niż LCO. Elastyczność projektowania pozwala inżynierom zmaksymalizować wydajność akumulatora w celu uzyskania optymalnej trwałości (żywności), maksymalnego prądu obciążenia (moc właściwa) lub dużej pojemności (energia właściwa). Na przykład wersja o długiej żywotności w ogniwie 18650 ma umiarkowaną pojemność wynoszącą zaledwie 1100 mAh; wersja o dużej pojemności to 1500 mAh. Czyste akumulatory litowo-manganowe nie są już dziś powszechne; można ich używać wyłącznie do zastosowań specjalnych.

Większość akumulatorów litowo-manganowych łączy się z tlenkiem litowo-niklowo-manganowo-kobaltowym (NMC), aby poprawić energię właściwą i przedłużyć żywotność. Ta kombinacja wydobywa to, co najlepsze z każdego systemu – LMO(NMC) jest wybierany do większości pojazdów elektrycznych, takich jak Nissan Leaf, Chevy Volt i BMW i3 (fotografia 4). Część LMO akumulatora, która może wynosić około 30 procent, zapewnia wysoki prąd podczas przyspieszania, część NMC zapewnia duży zasięg jazdy.

Badania nad akumulatorami litowo-jonowymi skupiają się głównie na łączeniu litowo-manganu z kobaltem, niklem, manganem i/lub aluminium jako aktywnym materiałem katody. W niektórych architekturach do anody dodaje się niewielką ilość krzemu. Zapewnia to 25-procentowy wzrost wydajności, jednak jest zwykle związane z krótszym cyklem życia ogniwa (krzem rośnie i kurczy się pod wpływem ładowania i rozładowywania, powodując naprężenia mechaniczne).

Akumulatory litowo-żelazowo-fosforanowe (LiFePO₄ – LFP)

Ogniwo LFP zapewnia dobrą wydajność elektrochemiczną przy niskiej rezystancji. Kluczowymi zaletami są:

- wysoki prąd znamionowy,
- długa żywotność, na poziomie 2000 cykli lub więcej,
- dobra stabilność termiczna,
- są uważane za jedne z najbezpieczniejszych,
- pozwalają na głębokie rozładowanie (nawet 100%), bez szkody dla ogniwa,
- spora tolerancja w przypadku przekroczenia parametrów.

Niższe napięcie nominalne wynoszące 3,2 V/ogniwo zmniejsza energię właściwą poniżej energii ogniwa z domieszką kobaltu. W przypadku



Fotografia 5. Klasyczny akumulator samochodowy z odniami LFP

większości akumulatorów niska temperatura zmniejsza wydajność, a podwyższona temperatura przechowywania skraca żywotność.

LFP charakteryzuje się większym samorozładowaniem niż inne akumulatory litowo-jonowe, co może powodować problemy z balanowaniem ogniw w miarę starzenia. Można temu zaradzić, kupując wysokiej jakości ogniwa i/lub stosując wyrafinowaną elektronikę sterującą, co zwiększa koszt pakietu. Czystość w procesie produkcji ma znaczenie dla długowieczności.

Akumulatory z ogniwami LFP często są używane do zastąpienia akumulatora kwasowo-ołowiowego (fotografia 5). Cztery ogniwa połączone szeregowo wytwarzają napięcie 12,80 V, czyli napięcie podobne do napięcia sześciu ogniw kwasowo-ołowiowych 2 V połączonych szeregowo. W przypadku czterech ogniw połączonych szeregowo, każde ogniwo osiąga maksymalne napięcie 3,60 V, co jest prawidłowym napięciem pełnego naładowania. LFP jest tolerancyjny na pewne przeładowanie, jednak utrzymywanie napięcia na poziomie 14,40 V przez dłuższy czas, jak ma to miejsce w przypadku większości pojazdów podczas długich podróży, może mieć niekorzystny wpływ. Niska temperatura zmniejsza wydajność akumulatora litowo-jonowego, co w skrajnych przypadkach może mieć wpływ na zdolność rozruchową.

Akumulatory z tytanem litu (Li₂TiO₃ – LTO)

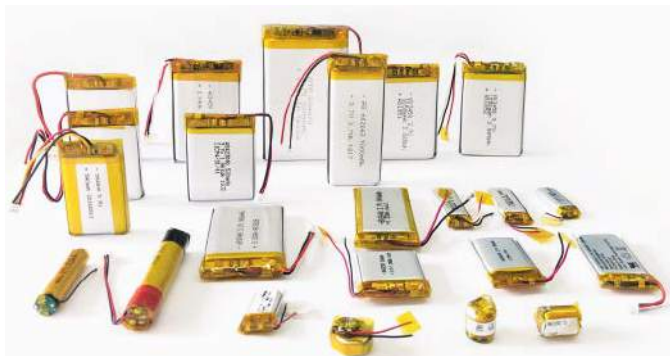
Baterie z anodami z tytanianu litu są znane od lat 80. XX wieku. LTO zastępuje grafit w anodzie typowego akumulatora litowo-jonowego, a materiał tworzy strukturę spinelową. Katodą może być tlenek litowo-manganowy lub NMC.

LTO ma nominalne napięcie ogniwa 2,40 V, można go szybko ładować i zapewnia wysoki prąd rozładowania wynoszący 10 C lub 10-krotność pojemności znamionowej. Mówi się, że liczba cykli jest wyższa niż w przypadku zwykłego akumulatora litowo-jonowego. LTO jest bezpieczny, ma doskonałe właściwości wyładowcze w niskich temperaturach i zachowuje 80% pojemności w temperaturze -30°C.

LTO ma przewagę nad konwencjonalnym akumulatorem litowo-jonowym z domieszką kobaltu i anodą grafitową, ponieważ charakteryzuje się zerowym naprężeniem i brakiem platerowania litem podczas



Fotografia 6. Ogniwo LTO typu 66160



Fotografia 7. Różne rodzaje akumulatorów litowo-polimerowych

szybkiego ładowania i ładowania w niskiej temperaturze. Stabilność termiczna w wysokiej temperaturze jest również lepsza niż w przypadku innych systemów litowo-jonowych, jednak bateria jest droga.

Energia właściwa wynosząca zaledwie 65 Wh/kg jest niska i porównywalna z energią ogniwi Ni-Cd. Ładuje się do 2,80 V/ogniwo, a koniec rozładowania wynosi 1,80 V/ogniwo. Typowe zastosowania to elektryczne układy napędowe, zasilacze UPS i oświetlenie uliczne zasilane energią słoneczną.

Akumulatory litowo-polimerowe (Li-Poly)

W akumulatorach litowych typu LCO, LMO, LMC stosowano i nadal stosuje się ciekłe elektrolity o różnym składzie, w tym wodne roztwory różnych bardziej i mniej bezpiecznych substancji zawierających lit. Istotnym wynalazkiem było zastąpienie cieczy stałym elektrolitem w postaci przewodzących polimerów, zawierających sole litu. Tak powstały akumulatory litowo-polimerowe, oznaczane Li-Po, LiPo lub LIP czyli akumulatory litowo-jonowe, zawierające kobalt lub mangan, gdzie ciekły elektrolit zastąpiono stałym.

Zasadniczo polimerowy elektrolit zwiększył rezystancję wewnętrzną, ale zwiększył bezpieczeństwo przez eliminację możliwości pożaru. Pierwsze akumulatory litowo-polimerowe były ogniwami cylindrycznymi, natomiast zastąpienie ciekłego elektrolitu stałym otworzyło drogę do realizacji akumulatorów o niemal dowolnym kształcie. Zaczęły się upowszechniać akumulatory litowe o kształcie prostokątnych poduszek (fotografia 7).

Tu warto wyjaśnić pewne niejasności. Otóż zasadniczo Li-Po lub LiPo oznacza akumulatory ze stałym, polimerowym elektrolitem. Jednak niedostateczna świadomość użytkowników i względy marketingowe spowodowały, że polimerowymi nazywa się też akumulatory z ciekłym elektrolitem, tylko mające kształt prostokątnych poduszek z tworzywa.

Dla niektórych wystarczającym powodem, by mówić o akumulatorach polimerowych była też obecność w środku separatora (porowatej folii) z tworzywa sztucznego. Ponieważ wcześniejsze wersje miały metalową obudowę, dla innych takim powodem była obecność miękkiej, plastikowej obudowy.

Pomimo tego rodzaju niejasności podstawowe zasady są proste: akumulatory polimerowe to też akumulatory litowo-jonowe, zwykle zawierające kobalt lub mangan, choć istnieją też „poduszkowe” LiFePO_4 i ich podstawowe parametry są takie, jak klasycznych z ciekłym elektrolitem.

Podsumowanie

Informacje podane w artykule sygnalizują najważniejsze zagadnienia, ale nie wyczerpują tematu akumulatorów litowych. Sytuacja rynkowa szybko się zmienia. Akumulatory i przechowywana w nich energia elektryczna stanowią siłę napędową nowoczesnych pojazdów elektrycznych, podobnie jak benzyna w samochodach z XX wieku. Producenci samochodów inwestują ogromne fundusze w rozwój tej branży.

Damian Sosnowski, EP

REKLAMA

przejrzyj i kup na
www.ulubionykiosk.pl

Banki energii

Zarządzanie ciepłem i ochrona przed warunkami środowiskowymi

Duża liczba instalacji fotowoltaicznych dołączonych do nieprzygotowanej na to sieci energetycznej wywołuje szkodliwe dla urządzeń podwyższenia napięcia. Zjawisko to występuje przy dużym nasłonecznieniu i powoduje również wyłączanie falowników, uniemożliwiając w ten sposób oddawanie energii i straty finansowe. Lekarstwem na te problemy mogą być banki energii, które pozwalają na oddawanie energii w szczycie energetycznym, co dodatkowo może poprawić parametry sieci i zwiększyć opłacalność inwestycji.

Bank energii to urządzenie zawierające ogniwa litowo-jonowe, które mają bardzo dużą pojemność i sprawność, wymagają jednak zapewnienia optymalnych warunków pracy. Dotyczy to przede wszystkim temperatury pracy, ale również innych warunków środowiskowych, takich jak wpływ wilgoci czy zanieczyszczeń chemicznych.

Systemy chłodzenia

Mimo wysokiej sprawności przetwarzania, mieszczącej się w granicach 96...99%, dużym problemem jest ciepło generowane w ogniwach zarówno w cyklu ładowania, jak i rozładowania. Z uwagi na to konieczne jest zainstalowanie efektywnego systemu chłodzenia – może to być chłodzenie pasywne lub aktywne. Z dużych różnic w radialnej i aksjalnej przewodności cieplnej, które dla ogniwa cylindrycznego wynoszą odpowiednio: 0,2...8 W/mK i 15...160 W/mK, wynika, że najefektywniejsze i jednocześnie najłatwiejsze jest chłodzenie dolnej powierzchni ogniwa.

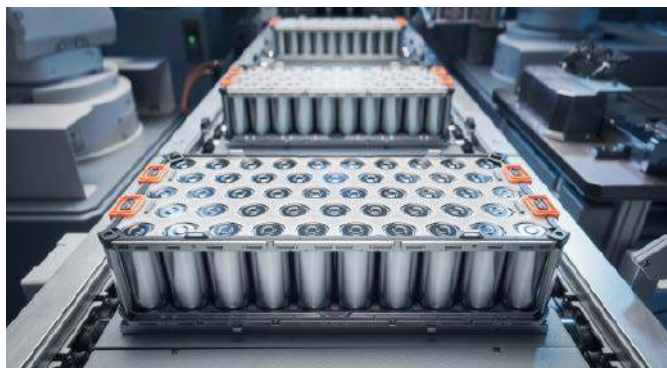
W przypadku niewielkich baterii często wystarcza pasywny radiator, na którym umieszczona jest termoprzewodząca podkładka o przewodności cieplnej od 1,5 W/mK do 3 W/mK. Parametry takie zapewniają podkładki typu Gap filler (**fotografia 1**) firmy Hala Contec z serii TGF m.in.:

- TGF-JXS-SI,
- TGF-M-SI (bez wzmocnienia włóknem szklanym),
- TGF-DXS-SI-GF,
- TGF-MXS-SI-GF (wzmocnione włóknem szklanym),
- TGF-R-NS (w aplikacjach bez silikonu).

Znakomitą alternatywą w przypadku wielkoseryjnej produkcji są dozowane materiały termoprzewodzące. Dzięki lepszemu wypełnieniu szczelin preparaty dozowane zapewniają efektywniejsze chłodzenie, a jednocześnie dużą powtarzalność i stabilność połączeń.



Fotografia 1. Gap fillery firmy HALA Contec



Więcej informacji:

Blelektronik

tel. 12 35 76 378, 696 483 020

kontakt@blelektronik.com.pl

www.blelektronik.com.pl



Dzięki niższemu kosztowi materiału inwestycje w urządzenia dozujące mogą szybko się zwrócić.

Duże zestawy baterii stosowane w przemysłowych bankach energii wymagają chłodzenia wymuszonego za pomocą wentylatorów, a często podobnie jak w samochodach elektrycznych za pomocą chłodnicy. Stosowane tam folie takie jak np. TFO-X-SI, TFO-Q-SI muszą mieć lepsze parametry termiczne. Ich przewodność wynosi około 5 W/mK.

W zależności od wymagań aplikacyjnych można dodatkowo zabezpieczyć baterie żywicą lub pianką silikonową np. SilsoLite (**fotografia 2**). Ma to na celu zmniejszenie różnic temperatur między poszczególnymi ogniwami. Zalewy takie poprawiają też znacząco odporność mechaniczną.

Wszystkie ww. materiały tworzą też bardzo odporną barierę izolacyjną o wytrzymałości powyżej 7 kV/mm, co umożliwia ich zastosowanie w roli izolacji funkcjonalnej oraz podstawowej. Możliwe jest również wykonanie z ich pomocą izolacji wzmocnionej, ale należy pamiętać o zapewnieniu wymaganej przez normy, powtarzalnej grubości warstwy.



Fotografia 2. Pianka silikonowa Silso Lite 21025

Ochrona przed uszkodzeniem

Ze względów bezpieczeństwa bardzo istotna jest ochrona przed uderzeniami mechanicznymi oraz wibracjami. Wszystkie dopuszczone do sprzedaży ogniwa oraz akumulatory muszą zostać pozytywnie zweryfikowane na zgodność z normą EN IEC 62619:2022 (EN 62133-2:2017 w przypadku urządzeń przenośnych), a dopuszczenie do transportu wymaga spełnienia zaleceń standardu UN38.3. Z uwagi na bardzo wysoką odporność na temperaturę oraz niepalność materiały na bazie silikonu zapobiegają propagacji uszkodzeń w bateriach litowych.

Zakładane wieloletnie bezobsługowe użytkowanie banku energii wymusza zastosowania, obudów o szczelności co najmniej IP55, a przypadku baterii nawet wyższej. Mimo to może dochodzić do kondensacji pary wodnej na elementach. W związku z tym warto dodatkowo zabezpieczyć je za pomocą silikonowych zalew QSIL553 i lakierów np. ACC15, ACC17. Należy przy tym pamiętać o wymogu stosowania materiałów o klasie palności V0.

Battery Management System

Prawidłową pracę baterii zapewnia moduł elektroniczny BMS (*Battery Management System*) zawierający zazwyczaj klucze przełączające w postaci tranzystorów MOS, układy pomiarowe i nadzorujące. Ciepło generowane w kluczach oraz w układach balansujących energię w celach musi zostać wyprowadzone poza moduł baterii, tak aby nie zwiększać dodatkowo jego temperatury. W tym przypadku pomocne są przekładki termoprzewodzące w formie folii, ale również dwuskładnikowe materiały dozowane.

W odróżnieniu od pakietu baterii, moc w BMS jest generowana punktowo, co powoduje konieczność sięgnięcia po znacznie skuteczniejsze materiały o przewodności co najmniej 2 W/mK, takich jak TGF-JXS-SI-A1 czy dozowany TDG-T-SI-2C (**fotografia 3**). W niektórych przypadkach zaleca się całkowite zalanie modułu elektroniki żywicą poliuretanową np. 8612 (ex ST13) lub zalewą silikonową SE3000, co poza poprawą warunków termicznych zapobiega przedostaniu się wilgoci do wnętrza.

Część aplikacji o dużej mocy może wymagać jeszcze skuteczniejszego chłodzenia. Warto wtedy rozważyć zastosowanie układu chłodzenia Heat pipe, który dzięki wyjątkowo wysokiej przewodności cieplnej >50 000 W/mK zapewnia półprzewodnikom optymalne warunki termiczne. Ich szczególną cechą jest bezobsługowość oraz długa żywotność.

Podsumowanie

Zapewnienie bezawaryjnej i bezpiecznej pracy banku energii wymaga zastosowania szeregu nowoczesnych rozwiązań w zakresie układów chłodzenia, ochrony przed wilgocią, zanieczyszczeniami oraz uderzeniami mechanicznymi.

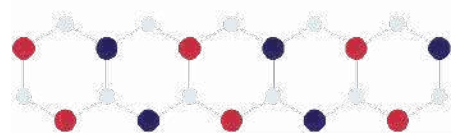
Blelektronik



Fotografia 3. Termoprzewodzący preparat dozowany



Optymalne rozwiązania
dla elektroniki i energetyki



Klejenie

Lakierowanie

Hermetyzacja

Odprowadzanie
ciepła



Robnor Resinlab®



Kisling HALA

FHU BL elektronik

ul. J. Conrada 65, 31-357 Kraków
tel. +48 12 357 63 78, +48 696 483 020
kontakt@blelektronik.com.pl
sklep.blelektronik.com.pl

e-sklep

Aplikacje enkoderów obrotowych

z zastosowaniem FPGA i układów SoC Microchip

W branży przemysłowej rośnie zapotrzebowanie na zaawansowane i niezawodne rozwiązania elektroniczne. Układy FPGA Smart Fusion 2 zapewniają inżynierom optymalne zasoby przy ściśle kontrolowanym, niskim poborze mocy i oferują jedno z najlepszych narzędzi do debugowania na rynku. Dzięki temu rozwiązują coraz bardziej krytyczne problemy i zapewniają wyjątkową niezawodność.

Zasada działania enkodera obrotowego

Enkodery obrotowe precyzyjnie mierzą położenie kątowe silnika elektrycznego, umożliwiając jego dokładne sterowanie. Enkodery te są zwykle montowane blisko wału silnika, tak jak pokazano na **rysunku 1**. Wiąże się to ze znacznymi ograniczeniami systemowymi:

- płytka drukowana (PCB) w kształcie pierścienia daje niewiele miejsca na rozmieszczenie komponentów,
- wysoka temperatura pracy silnika oznacza wysoką temperaturę elektroniki,
- wysoka temperatura oraz drgania i wibracje wpływają na zmniejszoną niezawodność i żywotność elektroniki.

Ograniczenia te powodują, że wszelkie rozwiązania elektroniczne wymagają najbardziej niezawodnych komponentów o niewielkim samonagrzewaniu i niewielkiej powierzchni. Ponieważ enkodery obrotowe są zwykle produkowane w dużych ilościach, koszt chłodzenia systemu i produkcji płytek PCB jest ważnym czynnikiem projektowym.

Dodatkowym czynnikiem istotnym w systemach tego typu jest bezpieczeństwo funkcjonalne, które zapewnia wiarygodność wartości pomiarowych i prawidłowe reakcje na stany awaryjne, takie jak doprowadzenie do stanu bezpiecznego – zatrzymanie.

Średnica typowej płytki PCB w kształcie pierścienia, przystosowanej do takiej aplikacji, spada do ok. 35 mm, co ogranicza fizyczny rozmiar wybranych komponentów do około 10×10 mm² lub mniej. Na płytę enkodera wpływa bezpośrednio ciepło wytwarzane przez silnik i w warunkach pracy może osiągnąć temperaturę 95°C, a okresowo nawet 105°C. W takich temperaturach samonagrzewanie się elektroniki na skutek prądów upływowych jest już bardzo istotne – niektóre półprzewodniki mogą nawet cierpieć na niekontrolowaną niestabilność cieplną.

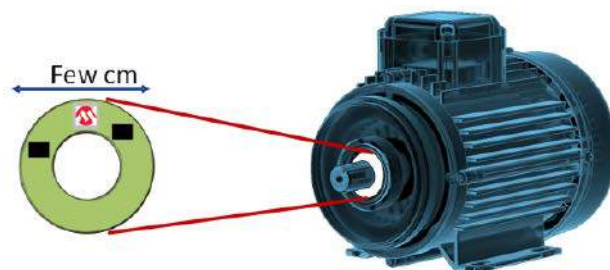
SmartFusion 2 SoC FPGA

Inżynierowie firmy Microchip specjalnie do takich aplikacji zaprojektowali układ SmartFusion 2 SoC FPGA w obudowie FCSG158. Jest to mikrosystem wyposażony w mikrokontroler Arm Cortex M3 (MCU) z dołączoną strukturą FPGA złożoną z około 25 000 elementów logicznych (LE)



i z pakietem bloków peryferyjnych o łącznych wymiarach 9×9 mm² (**rysunek 2**). Układ jest wysoce zoptymalizowany pod kątem trudnych warunków pracy i ograniczeń konstrukcyjnych.

Pomimo niewielkiej obudowy, układ został zoptymalizowany pod kątem routingu i umożliwia w pełni funkcjonalną aplikację za pomocą tylko dwóch warstw sygnału na płytce PCB. Dzięki temu płytka PCB nie jest skomplikowana i jest tania w produkcji. Wolna przestrzeń



Rysunek 1. Sposób zamontowania płytki PCB enkodera obrotowego w konstrukcji typowego silnika

More Resources in Low-Density Devices

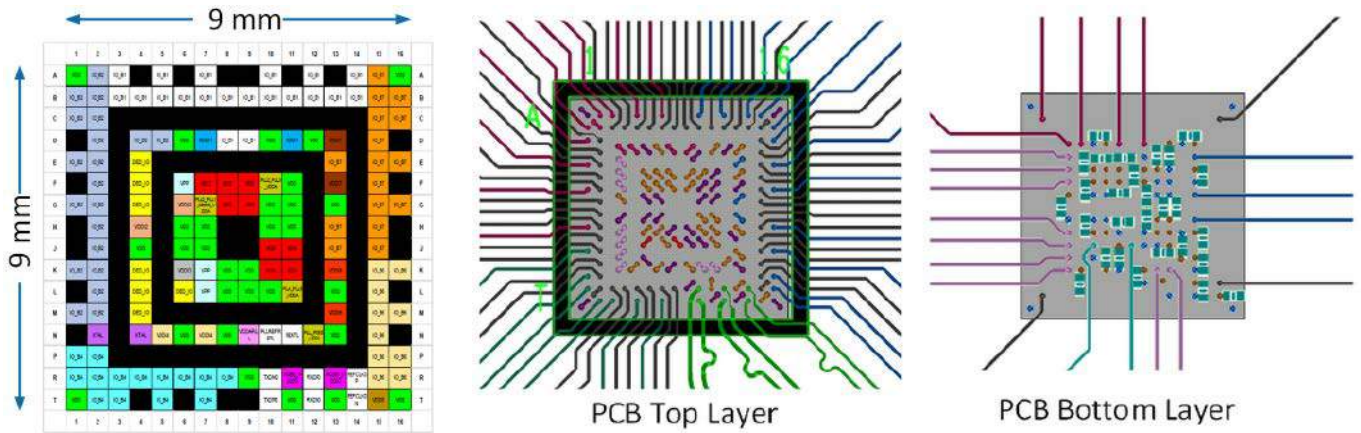
- Arm® Cortex®-M3 with embedded flash
- PCIe Gen2 support in 10K LE
- Comprehensive microcontroller subsystem

With Clear Advantages

- Lowest Power
 - Reduces total power by up to 50%
 - 70 mW per 5G SERDES (PCIe Gen2)
- Proven Security
 - Protection from overbuilding and cloning
 - Secure boot for FPGA and processors
- Exceptional Reliability
 - SEU immune zero FIT flash FPGA configuration
 - Reliable safety-critical and mission-critical systems



Rysunek 2. Układ SmartFusion 2 SoC FPGA



Rysunek 3. Układ wyprowadzeń SmartFusion 2 oraz schemat mozaiki ścieżek na 2-warstwowej płytce PCB

w siatce padów układu pozwala na zastosowanie ekonomicznych przelotek 0,3 mm dla pierścienia wewnętrznego i umieszczenie kondensatorów odsprężających bezpośrednio pod układem w pobliżu wyprowadzeń zasilających pośrodku pakietu (**rysunek 3**).

SmartFusion 2 FPGA zapewnia 82 wejścia/wyjścia, z których 70 może łączyć się z napięciem 3,3 V, a 12 z nich może natywnie łączyć się z napięciem 2,5 V lub poprzez dzielniki rezystancyjne z napięciem 3,3 V. Dodatkowo w sytuacjach wymagających szybkiej komunikacji dostępna jest jedna para transceiverów, która może pracować z szybkością do 5 Gb/s i obsługuje PCIe Gen2. Wbudowany w chip M3 transceiver można również zastosować do typowych zadań komunikacyjnych.

Dzięki zoptymalizowanej pod względem mocy architekturze układu FPGA SmartFusion 2 i wynikającemu z niej niewielkiemu samonagrzewaniu urządzenie może pracować w temperaturach otoczenia bliskich maksymalnie określonego zakresowi temperatur (100/125°C). Aby uzyskać dokładne szacunki zużycia energii i samonagrzewania urządzenia, zaleca się użycie narzędzia Microchip Power Estimator dla układów SmartFusion 2 SoC, IGLOO 2 FPGA lub PolarFire FPGA i SoC.

Bezpieczeństwo aplikacji

Jak wspomniano wcześniej, dane enkodera są często krytyczne dla bezpieczeństwa aplikacji, a projekt FPGA musi również mieć certyfikat bezpieczeństwa. Taki certyfikat bezpieczeństwa wspierany jest przez Microchip pakietem bezpieczeństwa SmartFusion 2/IGLOO 2 FPGA. Urządzenia SmartFusion 2 i IGLOO 2 FPGA mają certyfikat bezpieczeństwa funkcjonalnego zgodnie z normą IEC 61508. Odpowiedni pakiet bezpieczeństwa zawiera certyfikowane pod względem bezpieczeństwa środowisko programistyczne Libero SoC Design Suite wersja 18.3 SP4, 28 rdzeni IP, które są bardzo często wymagane

w projektach FPGA oraz podręcznik bezpieczeństwa dla tych urządzeń i narzędzi, aby obliczyć prawdopodobieństwo awarii sprzętu.

Ponieważ nasze układy FPGA i SoC są odporne na zakłócenia spowodowane pojedynczym zdarzeniem (SEU – *Single Event Upset*) w pamięci konfiguracyjnej, trwałe awarie sprzętu są jedynymi czynnikami wpływającymi na obliczenia FIT wymagane do dyskusji na temat bezpieczeństwa. FIT (*Failures in Time*) oznacza awarię w czasie, a jeden FIT oznacza jedną awarię w ciągu 109 godzin. W przypadku typowych układów FPGA bazujących na pamięci SRAM miękkie FIT powodowane przez SEU zależy od dokładnej architektury i złożoności wybranego urządzenia i zazwyczaj mieści się w zakresie około 400 FIT. Oznacza to, że funkcje istotne dla bezpieczeństwa zajmują zwykle tylko ułamek struktury FPGA. Do tego miękkiego FIT dodawane są trwałe awarie.

Wskaźnik FIT dla trwałych awarii ustala się na podstawie danych pomiarowych dostawcy przy podwyższonym poziomie obciążenia urządzenia, aby przyspieszyć efekty starzenia, które głównie przyczyniają się do awarii urządzeń. Wyniki tych pomiarów stanowią podstawę do obliczenia współczynnika awaryjności sprzętu w pożądanym warunkach operacyjnych w oparciu o zależność Arrheniusa:

$$\text{współczynnik awaryjności} = \frac{x^2 10^9}{2(A \cdot F \cdot \text{czas pracy})}$$

gdzie:

A.F – współczynnik przyspieszenia

Wizualny wynik tego obliczenia pokazuje wykres na **rysunku 4**.

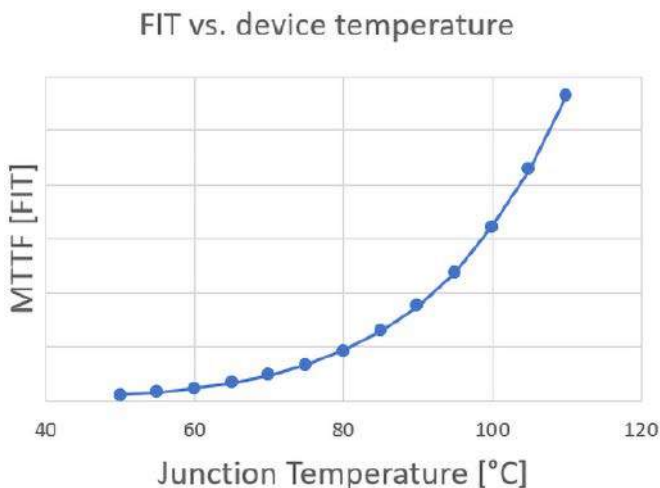
Podczas wykonywania obliczeń dla kompletnego układu FPGA uzyskany współczynnik FIT w temperaturach otoczenia > 90°C, w jakich często pracują enkodery obrotowe, może dać współczynnik awaryjności około 100 FIT. W przypadku konstrukcji bezpiecznych funkcjonalnie jest to zdecydowanie za dużo. Innymi słowy, należy albo utrzymywać niską temperaturę urządzenia, albo – ze względów bezpieczeństwa – zachować niewielki rozmiar struktury.

Nasz pakiet bezpieczeństwa zawiera odpowiednie dane i narzędzia do określenia współczynnika FIT części całego projektu, istotnej dla bezpieczeństwa. Pomaga to w uzyskaniu certyfikatu projektu. Ponadto współpracujemy z ekspertami branżowymi, którzy mogą świadczyć usługi doradcze w zakresie projektów bezpieczeństwa funkcjonalnego. W przypadku układów FPGA jest to Expleo, które jest dostępne w 30 różnych krajach na całym świecie.

Podsumowanie

Microchip zapewnia kompletny pakiet urządzeń zoptymalizowanych pod kątem zastosowań z enkoderami obrotowymi wraz z pakietem bezpieczeństwa i ekspertami doradczymi. Aby dowiedzieć się więcej, skontaktuj się z odpowiednim przedstawicielem firmy Microchip.

Martina Kellermann
Microchip



Rysunek 4. Obliczanie współczynnika awaryjności



Oscyloskop Siglent SDS1104X-U

Jeden z najlepszych budżetowych oscyloskopów cyfrowych

Oscyloskop, zwłaszcza cyfrowy, jest prawdopodobnie najbardziej użytecznym narzędziem w arsenale elektronika. Jest też narzędziem skomplikowanym, o wielu mniej lub bardziej istotnych parametrach czy funkcjach. Jest to inwestycja na lata, w dodatku często dość znaczna, dlatego warto dobrze przemyśleć wybór i za-inwestować w najlepsze dostępne narzędzie. Czy takim najlepszym narzędziem jest Siglent SDS1104X-U?

Na rynku nie brakuje marek i modeli oscyloskopów cyfrowych na każdą kieszeń. Często jednak te najtańsze mają „sztucznie zawyżone” parametry, kiepsko wykonane interfejsy i są źródłem błędów i frustracji, a nie użytecznymi narzędziami. Elektronikhobbysta może łatwo dokonać niewłaściwego wyrobu urządzenia „oscyloskopopodobnego” w cenie markowego urządzenia. Z kolei naprawdę dobre modele są bardzo kosztowne, a ich możliwości wykraczają daleko poza potrzeby większości, nawet zawodowych, elektroników. Na szczęście w ostatniej dekadzie nastąpił prawdziwy wysyp modeli budżetowych o przyzwoitych parametrach. W zaprezentowanej recenzji przyjrzymy się bliżej jednemu z nich, a ja postaram się wyjaśnić, dlaczego moim zdaniem jest to obecnie najlepszy oscyloskop cyfrowy dla hobbysty na polskim rynku.

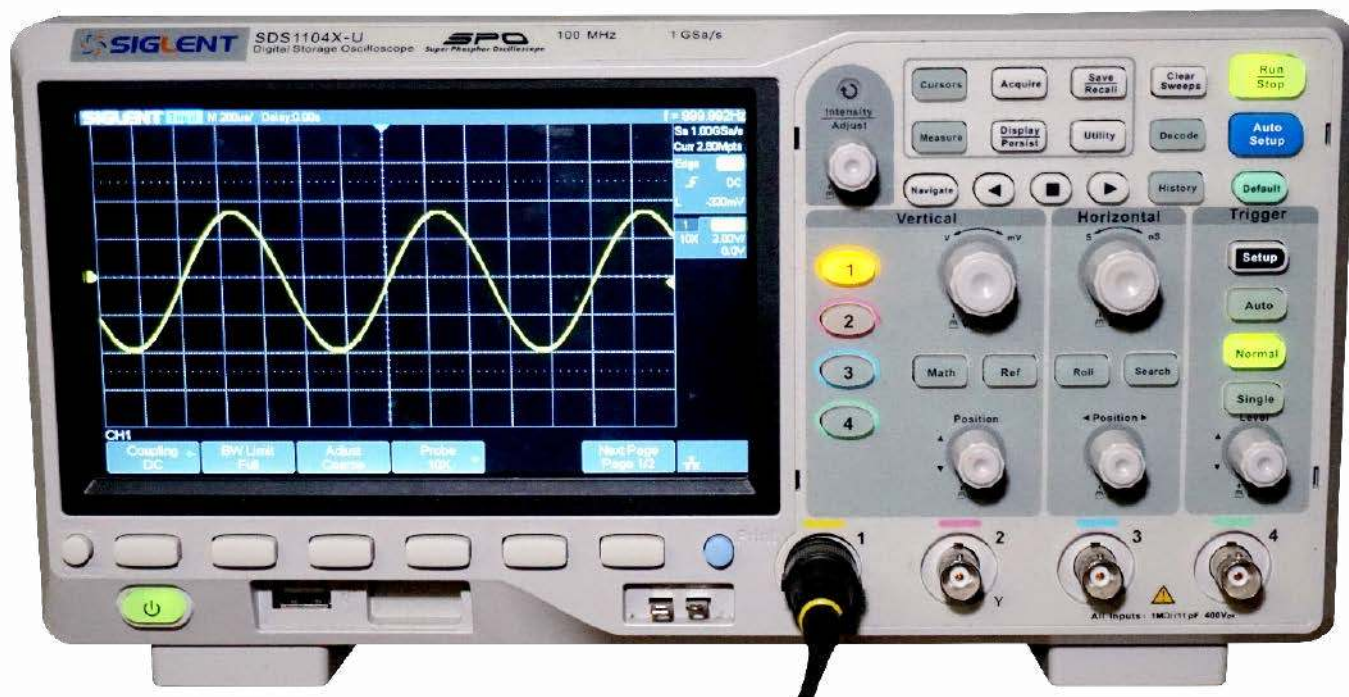
Podstawowe parametry oscyloskopu Siglent SDS1104X-U:

- 4 kanały o paśmie 100 MHz,
- częstotliwość próbkowania 1 GSa/s dla jednego kanału, 500 MSa/s dla dwóch i 250 MSa/s dla 3. i 4.,
- pamięć 14 mln próbek,
- wiele trybów wyzwalania i funkcji pomiarowych,
- dekodowanie popularnych protokołów szeregowych.

Zakup, unboxing i pierwsze wrażenia

Decyzję o zakupie oscyloskopu cyfrowego rozważałem dość długo, bo ponad rok. Rozważałem najpopularniejsze wśród amatorów marki: Owon, Hantek, Rigol i Siglent. Owon i Hantek oferuje przyzwoite oscyloskopy, ale wykonane niezbyt solidnie, co daje się odczuć zarówno w wykonaniu, jak i sposobie działania. Rigol wśród hobbystów jest najbardziej popularną marką, zwłaszcza modele DS1052E i DS1054Z. Główną przyczyną tego stanu rzeczy jest możliwość łatwego hakowania tych oscyloskopów po to, aby odblokować płatne opcje, jak dekodowanie protokołów, szersze pasmo czy większa pamięć próbek. Ja jednak nie chciałem wykonywać takich „modyfikacji” i wolałem wybrać oscyloskop, który ma te funkcje w standardzie. Chciałem też mieć model czterokanałowy.

Firma Siglent ma dwa podobne modele oferujące takie możliwości: SDS1104X-E oraz SDS1104X-U. Droższy model X-E wzbogacony jest



Fotografia 1. Sygnał sinusoidalny 1 kHz z generatora funkcyjnego na ekranie SDS1104X-U

o dodatkowe złącze, do którego można dołączyć dwie przystawki: generator sygnałów arbitralnych i analizator stanów logicznych. Pierwszy pozwala użyć oscyloskopu do rysowania wykresów Bodego, co bywa bardzo pomocne przy projektowaniu oraz testowaniu wzmacniaczy i filtrów. Analizator logiczny wbrew pozorom przydaje się raczej rzadko i tylko, gdy musimy korelować sygnały analogowe z danymi cyfrowymi na wielu liniach – osobiście takiej potrzeby nie odczułem ani razu.

Model SDS1104X-E ma też dwa przetworniki ADC, przez co użycie wszystkich czterech kanałów ogranicza częstotliwość próbkowania do 500 MSa/s. Siglent SDS1104X-U nie ma portu rozszerzeń (w tym miejscu obudowa ma mało gustowną zaślepkę), nie pozwala na generowanie wykresów Bodego i ma tylko jeden przetwornik ADC. W zamian jest tańszy od swojego starszego brata. W styczniu 2022 roku różnica cenowa między modelami wynosiła 280 zł, przy czym model X-U kosztował 1870 zł. W chwili pisania tych słów, czyli w październiku 2023 roku model X-U kosztuje już 2054 zł, a model X-E 2361 zł. – różnica wzrosła do 307 złotych.

Oscyloskop dotarł do mnie zapakowany w dość solidny, podwójny karton. Standardowo w zestawie były cztery sondy oscyloskopowe, kabel USB oraz kabel zasilający w standardzie IEC. Znajdziemy też skróconą instrukcję obsługi i certyfikat kalibracji. Sam sprzęt wykonany jest z tworzywa sztucznego, ale pod nim kryje się metalowa skorupa ekranująca.

Na panelu frontowym oscyloskop ma cztery gniazda BNC dla sond, są one pozbawione jakichkolwiek dodatkowych styków – w końcu to produkt budżetowy. Na lewo od nich znajdują się dwa styki do kompensacji sond, zaś obok włącznika zasilania gniazdo USB Host. Pierwsze dobre wrażenie psuje nieco wspomniana wyżej zaśleпка portu rozszerzeń między gniazdem USB a stykami do kompensacji. Z tyłu znajdziemy gniazda USB typu B, zasilania, gniazdo BNC Trigger Out/Pass-Fail, gniazdo Ethernet oraz otwór Kensington Lock.

Masa urządzenia jest niewielka, dlatego urządzenie można z łatwością przenosić i transportować. Przyciski są gumowe o wyraźnym wyczuwalnym kliknięciu, część jest podświetlona (fotografia tytułowa), co jest typowym rozwiązaniem. Dla mnie pewną wadą jest to, że tylko dwa enkodery mają wyczuwalny klik, są to pokręta regulacji wzmocnienia i podstawy czasu. Moim zdaniem pokręta uniwersalne też powinno klikać, co by ułatwiło ustawianie wartości licznych parametrów ukrytych w menu. Firma Siglent ma też pewien problem

z estetyką – na panelu znajdziemy kilka różnych czcionek, a pojedyncze przyciski mają do tego różne kolory – trudno to uzasadnić. Dlaczego na przykład przycisk Setup w sekcji wyzwalania ma czarny kolor, gdy wszystkie inne przyciski „funkcyjne” są szare lub białe?

Sensowne za to jest umiejscowienie przycisku Auto Setup zaraz pod podświetlonym przyciskiem Run/Stop. Po jego użyciu oscyloskop próbuje ustawić kluczowe parametry tak, by pokazać sygnały, do których podłączone są sondy. Trwa to około 10 sekund. Ciekawym dodatkiem jest umieszczony niżej przycisk Default, który przywraca ustawienia domyślne. Umiejscowienie jest mało fortunne, bo kilka razy go nacisnąłem przez przypadek, a przywraca on też ustawienia siatki na ekranie, co bywa irytujące. Na szczęście można zapisać bieżące ustawienia jako domyślne – funkcja ta przydaje się zwłaszcza na początku, gdy jeszcze uczymy się obsługi i coś namieszamy.

Ostatnia drobna sprawa to kolory diod LED przycisków wyboru kanału – kolory te nie zgadzają się z kolorami na obudowie, na ekranie czy nawet na pierścieniach sond oscyloskopowych. A lepsze ich dobranie nie wpłynęłoby na koszty produkcji.

Ekran TFT o przekątnej siedmiu cali i rozdzielczości 800×480 pikseli jest wystarczająco czytelny. Przeszkadzać mogą ograniczenia tej technologii, ale niewątpliwą zaletą będzie długowieczność matrycy. Interfejs jest relatywnie wygodny – przyciski funkcyjne są pod ekranem i odpowiadają im odpowiednie pola na ekranie. Fotografia 1 pokazuje przykładowy sygnał z generatora funkcyjnego na tym ekranie. Mamy też specjalny przycisk do zapisywania zrzutów ekranu, który możemy też skonfigurować w ustawieniach – poza kilkoma formatami obrazu możemy zachować też zebrane próbki w formie pliku CSV. W trakcie używania nie zauważyłem, by interfejs zwalniał z powodu dużej liczby aktywnych funkcji, a sam oscyloskop zawiesił się tylko jeden raz, gdy próbowałem zapisać zrzut ekranu natychmiast po podłączeniu pamięci USB. Jedynym minusem, jak dla mnie, jest dobór pewnych kolorów, wynika to tylko z mojej wady

Tabela 1. Parametry sond oscyloskopowych

Tłumienie	1×	10×
Pasma przenoszenia	6 MHz	100 MHz
Impedancja wejściowa	1 MΩ	10 MΩ
Pojemność wejściowa	85...120 pF	18...22 pF
Maks. napięcie wejściowe	<300 Vp-p	<600 Vp-p

wzroku – miejscami kontrast jest zwyczajnie za mały. Do interfejsu jeszcze wrócę przy okazji omawiania różnych funkcji.

Siglent dołączył do oscyloskopu cztery sondy własnej produkcji, model PP510. Na oficjalnej stronie www.siglent.eu sondy te kosztują 16 euro. Są to standardowe sondy 1X/10X o paśmie 100 MHz, dokładne parametry pokazuje tabela 1.

Poza samymi sondami dołączone są dodatkowe akcesoria, takie jak nakładki pozwalające na wygodniejszy pomiar między elementami, plastikowe pierścienie w czterech kolorach do oznaczania kanałów – pasują do barw na oscyloskopie, wkrętaki do kalibracji czy sprężyste styki GND przydatne przy pomiarach w.c.z. Jakościowo sondy wykonane są lepiej, niż można się spodziewać, biorąc pod uwagę cenę, ale nadal są to produkty budżetowe.

Podstawowe parametry

Siglent, chcąc odbić rynek hobbystów z rąk Rigola, zdecydował się opracować oscyloskop, który będzie oferował to samo lub więcej niż Rigol DS1054Z. Gdy decydowałem się na zakup, nawet model SDS1104X-E był tańszy od DS1054Z, nie wspominając już o omawianym tutaj modelu X-U. Obecnie cena produktu Rigola jest zbliżona do SDS1104X-U, głównie dlatego, że Rigol wypuścił na rynek nowe modele serii DHO800, co wpłynęło na spadek cen starszych modeli. Model Siglenta oferuje cztery kanały o paśmie 100 MHz, z częstotliwością próbkowania 1 GSa/s dla jednego kanału, 500 MSa/s dla dwóch i 250 MSa/s dla trzech i czterech. Pamięć akwizycji ma tylko 14 milionów punktów i to bez możliwości rozszerzenia, ale częstotliwość odświeżania wynosi 100 000 przebiegów na sekundę w trybie normalnym i 400 000 przebiegów na sekundę w trybie sekwencyjnym. Zostało to potwierdzone przez obserwowanie przebiegu na wyjściu Trigger Out.

SDS1104X-U używa pojedynczego przetwornika ADC HMC1511 od Analog Devices. Przetwornik ten oferuje cztery różnicowe kanały analogowe, rozdzielczość 8 bitów, wewnętrzne napięcie odniesienia 1 V oraz elastyczny system taktowania przetworników. Każdy kanał ma własny programowalny wzmacniacz. Rozdzielczość pionowa oscyloskopu wynosi od 1 mV na działkę do 10 V na działkę w skali 1-2-5. Dla każdego kanału można dodatkowo ustawić podział podłączonej sondy, dzięki czemu wartości wyświetlane na ekranie są podawane prawidłowo. Można też przełączać jednostki, jeśli podłączymy na przykład sondę prądową.

Każdy kanał ma też własny ogranicznik pasma do 20 MHz. Minimalne pasmo przenoszenia przy sprzężeniu AC jest poniżej 2 Hz. Pod względem zniekształceń wzmocnienia Siglent SDS1104X-U może się pochwalić liniowością tegoż na poziomie poniżej ± 1 dB do 10 MHz, poniżej ± 2 dB dla zakresu 10...50 MHz oraz około $+2$ dB od 50 MHz do 100 MHz. W testach różnych użytkowników pasmo -3 dB sięga często powyżej 120 MHz, ale to mocno zależy od egzemplarza – Siglent gwarantuje pasmo do 100 MHz. W praktyce amatorskiej nieczęsto zdarza się pracować przy tak wysokich częstotliwościach – mnie w ostatniej dekadzie zdarzyło się to tylko dwa razy.

Podstawa czasu ma zakres od 2 ns do 100 s, ponownie w skali 1-2-5. Przesunięcie fazowe między kanałami wynosi mniej niż 100 ps, zaś sama podstawa czasu ma dokładność <25 ppm. Są to wartości typowe dla instrumentu tej klasy. Warto pamiętać, że wraz z wiekiem dokładność podstawy będzie maleć ze względu na starzenie się generatora. Przy okazji oscyloskop ma sprzętowy licznik częstotliwości. Wpływ na dokładność poziomą i pionową ma też temperatura pracy, dlatego Siglent zaleca, aby przy zmianie

warunków, w jakich znajduje się oscyloskop, dokonać automatycznej kalibracji.

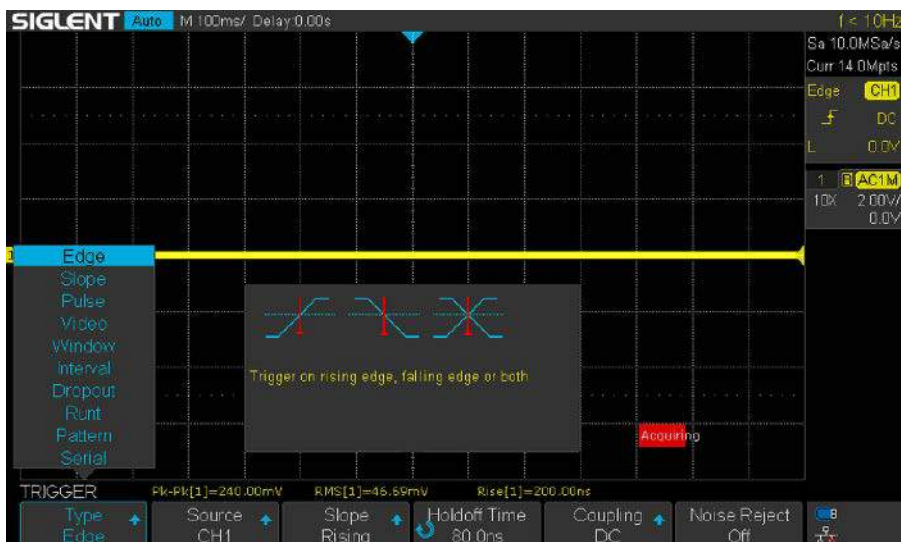
Wyzwalanie, akwizycja i opcje wyświetlania

Siglent SDS1104X-U ma dość rozbudowane możliwości wyzwalania. Możemy wybrać jeden z czterech kanałów, ale możemy użyć też linii zasilania oscyloskopu, co pozwala na łatwą synchronizację do napięcia sieci elektrycznej. SDS1104X-U oferuje trzy standardowe tryby wyzwalania: Auto, Normal i Single. W trybie Auto, jeśli warunki wyzwalania nie zostaną spełnione, oscyloskop i tak dokona akwizycji, byle coś pokazać na ekranie. W trybie Normal akwizycja nastąpi za każdym razem, gdy warunki wyzwalania zostaną spełnione, ale tylko wtedy. Tryb Single oczywiście dokona pojedynczej akwizycji, gdy warunki zostaną spełnione, po czym wyświetli wynik na ekranie. W przeciwieństwie do DS1054Z Rigola Siglent zdecydował się udostępnić wszystkie typy wyzwalania w cenie produktu, co było jednym z powodów, dla którego wybrałem właśnie ich oscyloskop. Menu wyboru typów wyzwalania (rysunek 1) oferuje od razu odpowiedź wyjaśniającą ich działanie.

Domyślnym typem wyzwalania jest wyzwalanie zboczem narastającym, ale możemy też wyzwalać zboczem opadającym czy oboma zboczami. Typ Slope pozwala wyzwalać oscyloskop, gdy przejście między dwoma punktami na zboczu narastającym lub opadającym trwa mniej lub więcej czasu, niż wynosi zadana wartość. Z kolei opcja Pulse Width oscyloskop jest wyzwalany, gdy stan wysoki lub niski sygnału PWM trwa krócej lub dłużej, niż wynosi zadana wartość. Ta opcja na pewno może się przydać przy badaniu przetwornic, falowników i innych układów używających PWM.

Typ wyzwalania Window reaguje, gdy przebieg znajdzie się między dwoma zadanymi progami, przypomina to tryb Slope, ale bez składnika czasowego. Interval zaś wyzwalą wtedy, gdy dwa sąsiadujące zbocza narastające lub opadające są od siebie oddalone w czasie o mniej lub więcej, niż wynosi zadana wartość. Wyzwolenie nastąpi przy tym drugim zboczu. Podobnie działa Dropout – tu wyzwolenie nastąpi, gdy następne zbocze narastające lub opadające nie pojawi się po upływie zadanego czasu, lub gdy poziom sygnału nie zmieni się przez zadany czas. Ostatnim prostym typem wyzwalania jest Runt – wyzwolenie nastąpi wtedy, gdy sygnał przekroczy jeden z wybranych poziomów wyzwalania, ale nie przekroczy drugiego – pozwala to zaobserwować sytuację, gdy z jakiegoś powodu sygnał nie osiąga pożądanej amplitudy w losowych momentach.

SDS1104X-U oferuje też bardziej złożone typy wyzwalania. Pattern wykonuje operacje logiczne na dwóch kanałach, wyzwolenie następuje, gdy spełniony zostanie wybrany warunek. Do wyboru są operacje AND, OR, NAND i NOR. Wyzwolenie nastąpi, gdy zostanie spełnione



Rysunek 1. Menu wyborów typów wyzwalania

zarówno kryterium czasowe, a jednocześnie wynik operacji logicznej będzie fałszywy. Dla przykładu można ustawić warunek AND, dla kanału 1 stan wysoki, dla kanału 2 zaś stan niski. Gdy oba kanały będą miały stan wysoki w tym samym czasie, nastąpi wyzwolenie.

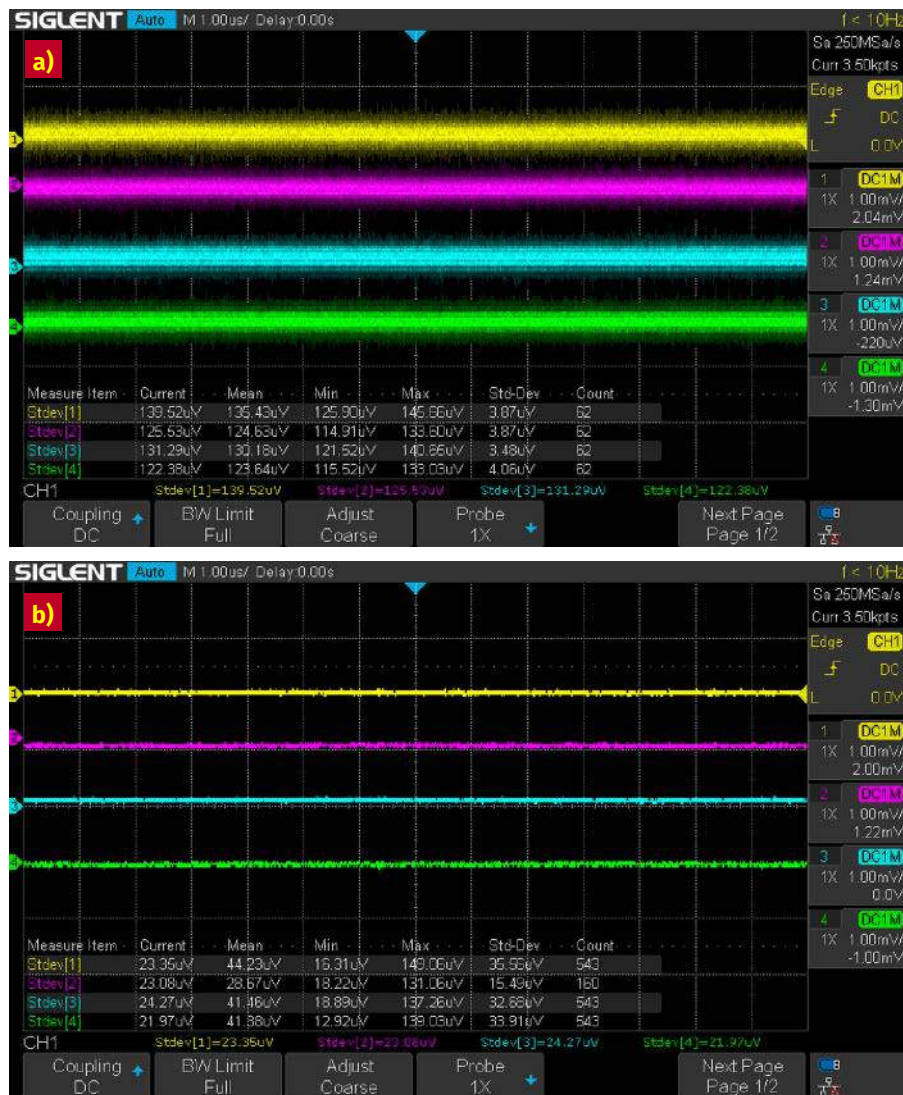
Opcją, która nie będzie zbyt często używana, jest wyzwalanie sygnałem wideo. Oscyloskop rozpoznaje standardy PAL i NTSC oraz HDTV 720p/1080p/720i/1080i z częstotliwościami odświeżania 50/60 klatek, ale też można wybrać własne ustawienia. Można nawet wybrać linię obrazu, na której oscyloskop zostanie wyzwolony.

Ostatnim, ale zarazem najbardziej użytecznym, typem wyzwalania jest wyzwalanie transmisją szeregową. SDS1104X-U pozwala wyzwać i dekodować protokoły I²C, SPI, UART, CAN oraz LIN. Funkcja ta jest bardzo rozbudowana i pozwala na dogłębną diagnostykę problemów komunikacyjnych. Oscyloskop potrafi rozpoznać różne stany, jakie mogą występować w tych protokołach i na ich podstawie się wyzwać. Dla przykładu, badając interfejs I²C, możemy wyzwać oscyloskop za każdym razem, gdy nastąpi stan No Ack. Możemy też ustawić określony adres lub określone dane, by oscyloskop reagował tylko na nie lub na ich brak. Te opcje doceni każdy, kto kiedykolwiek próbował uruchomić taką łączność, ucząc się programowania mikrokontrolerów.

Warto też pamiętać o tym, że system wyzwalania ma dodatkowe możliwości, takie jak sprzężenie AD/AC czy LF/HF Rejection, czyli filtrowanie sygnałów poniżej 2 MHz (LF) lub powyżej 1,2 MHz (HF). Jest też oddzielny filtr szumów. Te funkcje są niezależne od ustawień kanałów.

System akwizycji i wyświetlania informacji w SDS1104X-U jest również dość rozbudowany i oferuje możliwości podobne, co urządzenia konkurencji. Mamy więc tryb wyświetlania X-T, X-Y oraz Roll – ten ostatni emuluje zachowanie aparatu rysującego wykres na wstędze przesuwającego się papieru. Siglent wzbogacił swoje oscyloskopy o całkiem dobrze działającą emulację zachowania oscyloskopów analogowych, nazywaną SPO, czyli Super Phosphor. Oferuje on 256-poziomowy gradient intensywności oraz możliwość użycia gradientu kolorów w stylu kamer termowizyjnych. Dodatkowo możemy jeszcze użyć funkcji Persistence, dzięki czemu bufor obrazu wypełniany jest kolejnymi akwizycjami, a wcześniejsze są wymazywane po upływie zadanego czasu (1, 5, 10 lub 30 sekund) lub są zachowywane przez nieskończenie długi czas (lub do wyczyszczenia ekranu). Dzięki SPO i funkcji Persistence oraz wysokiej częstotliwości akwizycji łatwiej jest uchwycić na ekranie zdarzenia rzadkie, jak zniekształcone sygnały. Potem można zastosować odpowiedni do sytuacji typ wyzwalania, by dokładniej obejrzeć problematyczny sygnał. Siglent pozwala wyświetlać te ślady w formie wektorów lub punktów, z interpolacją $\sin(x)/x$ i liniową.

Siglent SDS1104X-U ma cztery tryby akwizycji: normalny, Peak Detect, Average i Eres. W trybie normalnym pozyskiwane są próbki w równych odstępach czasu, a potem od razu są wyświetlane. Jest to domyślny tryb pracy i sprawdza się w większości sytuacji. W trybie Peak Detect oscyloskop stara się „wyłapać” najwyższe i najniższe wartości napięcia w każdym interwale próbkowania. Dzięki temu może wyświetlić sygnały zbyt szybkie dla trybu normalnego. Siglent w instrukcji



Rysunek 2. Szumy wejściowe wszystkich kanałów z funkcją uśredniania wyłączoną (a) i włączoną (b)

pokazuje przykład sygnału PWM o częstotliwości 1 kHz i wypełnieniu 0,1%, który w trybie normalnym jest praktycznie niewidzialny, choć wyzwalanie nastąpiło właśnie na takiej „szpilce”. Wadą tego trybu jest zwiększony poziom szumu w sygnale. Tryb Average uśrednia wartości z kolejnych akwizycji sygnału celem uśrednienia poziomu szumu. Domyślnie uśrednia 16 kolejnych akwizycji, ale można ustawić nawet 1024 akwizycje do uśrednienia. Odświeżanie mocno zwalnia w takim wypadku, gdyż wymaga zebrania odpowiedniej liczby akwizycji, by je uśrednić. Ponadto uśrednianie tego typu może zniekształcić wyświetlany sygnał. Dla przykładu sygnał z modulacją FM przy niższych podstawach czasu w trybie Average nie wygląda wcale jak sygnał FM. Przykład użycia uśredniania można zademonstrować, mierząc szumy własne wejść oscyloskopu. Na rysunku 2a widzimy szumy w trybie normalnym, a na rysunku 2b w trybie uśredniania. Ostatni tryb, Eres, też dokonuje uśredniania, ale nie uśrednia próbek z kolejnych akwizycji, lecz uśrednia sąsiednie próbki w bieżącej akwizycji. W efekcie szum wysokiej częstotliwości zostaje uśredniony, rozdzielczość pionowa rośnie, nawet do 11 bitów, ale pasmo przenoszenia ulega silnej redukcji. Za to tryb ten działa w czasie rzeczywistym oraz w trybie Single Shot.

Ostatnią, bardzo użyteczną funkcją jest sekwencyjny tryb akwizycji. Zwiększa on częstotliwość akwizycji do nawet 400 tysięcy przebiegów na sekundę, lecz nie działa w czasie rzeczywistym – na wynik musimy poczekać. Jak to działa? Za każdym wyzwoleniem do pamięci zapisywany jest rekord o określonej długości, aż pamięć zostanie wypełniona. Potem możemy przeglądać kolejne akwizycje jedna po drugiej.

Tabela 2. Pełna specyfikacja oscyloskopu Siglent SDS1104X-U

Parametr	Wartość
Pasma przenoszenia	100 MHz
Częstotliwość próbkowania	Częstotliwość próbkowania 1 GSa/s (1 kanał), 500 MSa/s (2 kanały), 250 MSa/s (4 kanały)
Pamięć akwizycji	14 Mpts
Częstotliwość odświeżania przebiegu	100 wfm/s (tryb normalny), 400 wfm/s (tryb sekwencyjny)
Typ wyzwiania	Edge, Slope, Pulse Width, Window, Runt, Interval, Dropout, Pattern, Video, Serial
Rozdzielczość przetwornika ADC	8 bitów
Skala pionowa (sonda 1X)	1 mV/dz...10 V/dz.
Podstawa czasu	2 ns/dz...100 s/dz.
Stopnie intensywności (SPO)	256 poziomów
Tryby wyświetlania	X-T, X-Y, Roll
Dokładność podstawy czasu	±25 ppm
Sondy	4×Siglent PP510
Wyświetlacz	LCD TFT 7" 800×480
Waga	2,6 kg (bez opakowania)

Do czego to może się przydać? Na przykład do obserwowania szybkich zdarzeń, które zdarzają się w długich odstępach czasu. Skrajnym przykładem może być chęć odczytu danych w interfejsie SPI o prędkości 10 Mbps, gdzie transmisje następują raz na pół sekundy. W normalnym trybie pracy możemy albo zaobserwować pojedynczy pakiet danych, albo (po zwiększeniu podstawy do 500 ms...2 s na działkę) że „coś” się dzieje. Gdy używamy dwóch kanałów przy 500 ms na działkę, częstotliwość próbkowania wynosi 1 MSa/s, a głębokość pamięci tylko 7 Mpts. W trybie sekwencyjnym częstotliwość próbkowania może być dużo wyższa, długość każdego rekordu zaś dobrana do długości pakietu. W efekcie zbierzemy same pakiety danych bez „pustych” przestrzeni między nimi.

Funkcje pomiarowe, matematyczne i inne

SDS1104X-U, jak każdy oscyloskop cyfrowy, pozwala na pomiary badanych sygnałów. Mamy całkiem bogaty zestaw opcji pomiarowych, bo aż 38 różnych pozycji. Menu wyboru pomiarów prezentuje **rysunek 3**. Na **rysunku 4** został pokazany przykład pomiaru amplitudy (pk-pk i RMS) oraz czasu narastania sygnału prostokątnego. W praktyce jednak ich przydatność bywa różna. Większość pomiarów dotyczy pojedynczego kanału, ale jest grupa dziesięciu, które porównują dwa kanały ze sobą. Oczywiście mamy też dostępne pomiary kursorami (**rysunek 5**) oraz bardziej rozbudowaną statystykę. Mój wcześniejszy oscyloskop był modelem analogowym, więc choć funkcje pomiaru są standardem w świecie oscyloskopów cyfrowych, to dla mnie skok jakościowy jest nieporównywalny. Nie dość, że wszystkie potrzebne dane są czytelnie prezentowane, to dokładność jest definitywnie większa niż ręczne liczenie działek na niezbyt dużej lampie oscyloskopowej, jaką ma mój stary radziecki instrument.

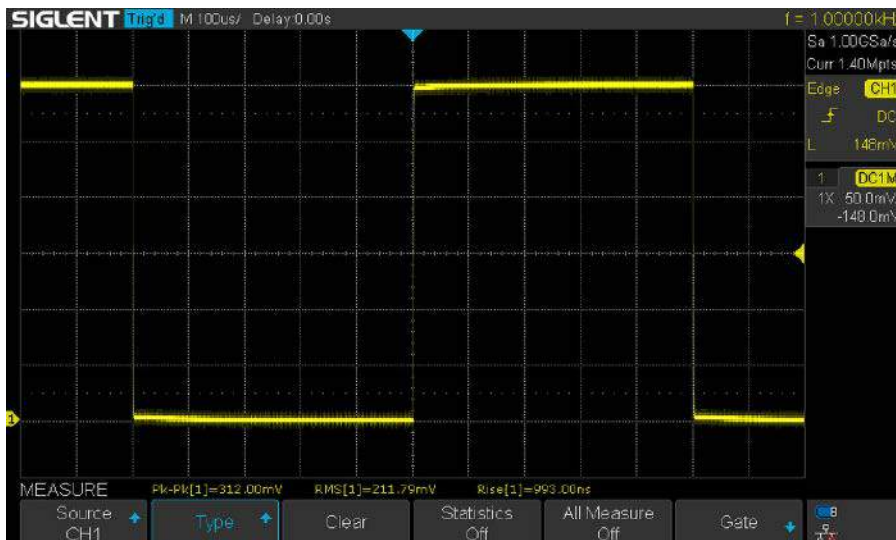
Przy okazji pomiarów możemy w dość prosty sposób sprawdzić poziom szumów własnych

wejść Siglenta SDS1104X-U. Co więcej, tą samą metodą możemy porównać wyniki każdego oscyloskopu cyfrowego. Procedura wygląda następująco:

1. Ustawiamy podstawę czasu na 1 μ s/dz. Każdy kanał ustawiamy na 1 mV/dz ze sprzężeniem DC;



Rysunek 3. Menu wyboru pomiarów. Każdy pomiar ma ilustrację i krótki opis



Rysunek 4. Pomiary sygnału testowego 1 kHz

2. Z menu measure wybieramy opcję Standard Deviation i włączamy wyświetlanie statystyk;
3. Notujemy wartość średnią (Mean) dla każdego kanału z osobna i dla różnych kombinacji kanałów. Te wartości pokazują poziom szumów RMS.

Rezultaty dla mojego egzemplarza Siglent SDS1104X-U zestawilem w tabeli 3.

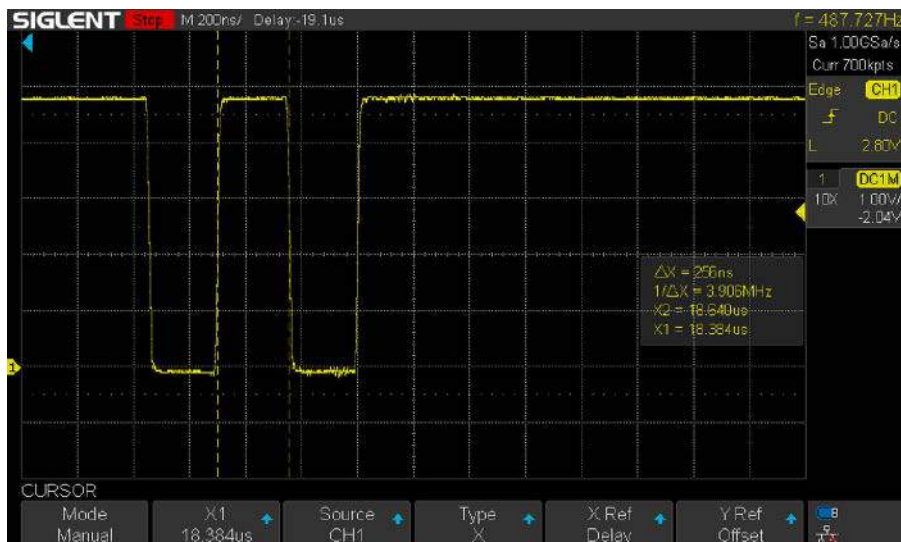
Siglent dodał też standardowe operacje matematyczne, które można wykonać na jednym lub dwóch kanałach. Wynik operacji jest pokazany jako dodatkowy biały ślad na ekranie. Możemy zatem obliczać sumę i różnicę dwóch kanałów, a także mnożyć jeden przez drugi lub je dzielić. Wynik może mieć jedną z kilku różnych jednostek. Do czego może się to przydać? Możemy na przykład mierzyć napięcie po obu stronach rezystora włączonego w szereg z linią zasilania układu, i odejmując jedną wartość od drugiej, uzyskać wykres zmian prądu. Kolejne dwie operacje to d/dt (differentiate), czyli różniczkowanie, oraz dt (integrate), czyli całkowanie. Operacja różniczkowania pozwala zbadać, jak szybko zmienia się napięcie w jednostce czasu, czyli slew rate. Przydaje się to do badania wzmacniaczy operacyjnych. Z kolei całka pozwala nam mierzyć energię impulsu (w woltach na sekundę) lub obszar pod wykresem. Trzecia operacja to pierwiastek kwadratowy sygnału.

Ostatnią operacją matematyczną dostępną w tym oscyloskopie jest funkcja FFT, czyli szybka transformata Fouriera. Zmienia to oscyloskop w analizator widma. Możemy zobaczyć zawartość harmonicznych i zniekształceń w badanym sygnale, poziom szumów w zasilaczu prądu stałego czy, jak sugeruje Siglent w instrukcji obsługi, badać wibracje. Niestety, ta sama instrukcja nie wyjaśnia zbyt dokładnie, jak używać funkcji FFT. W testach porównawczych między różnymi oscyloskopami Siglent wypada całkiem dobrze i działa dość szybko. Sam nie mam zbyt dużego doświadczenia w używaniu funkcji FFT i uważam, że Siglent dobrze by zrobił, dodając dokładniejsze wyjaśnienie, jak z niej korzystać.

Ciekawą funkcją, o której niewielu użytkowników oscyloskopów wie, jest opcja zapisu przebiegów referencyjnych. SDS1104X-U może zapisać i wyświetlić do czterech takich przebiegów naraz. Dwa typowe zastosowania to chęć porównania sygnału przed i po zmianie wartości jednego z elementów badanego obwodu. Dokonujemy pomiaru przed zmianą, po czym w menu REF zapisujemy i wyświetlamy ślad. Potem możemy zmienić wybrany komponent i ponownie uchwycić sygnał. Inną opcją jest uzyskanie dodatkowego (wirtualnego) kanału pod warunkiem, że źródłem wyzwalania będzie zawsze to samo zdarzenie. W tym miejscu objawia się kolejny raz problem kontrastu – kolorem przebiegu referencyjnego jest ciemnoniebieski, w dodatku jest nałożony na przebieg normalny i go zasłania. Jest to jednak istotny błąd interfejsu, który powinien być poprawiony.

Podsumowanie, czyli SDS1104X-U w praktyce

Oscyloskop mam od ponad półtora roku. Poza drobnymi problemami interfejsu nic mu nie można zarzucić. Producent odrobił lekcję i przygotował urządzenie lepsze niż oferta Rigola, który do tej pory rządził w segmencie oscyloskopów dla hobbystów i budżetowych zawodowców. Siglent zaoferował w standardzie to, co Rigol daje jako płatne opcje (wychodząc z założenia, że hobbysci i tak te opcje sobie odblokują, hakując urządzenia, ale firmy już uczciwie zapłacą), co moim



Rysunek 5. Pomiary czasu trwania stanu wysokiego za pomocą kursorów

Tabela 3. Poziom szumów własnych wejść oscyloskopu

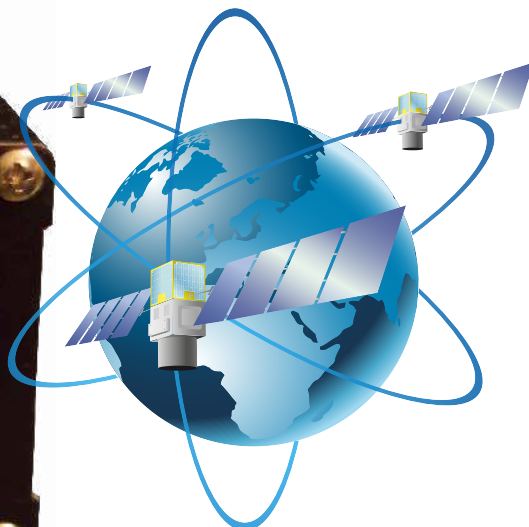
Kanał 1	Kanał 2	Kanał 3	Kanał 4
87,73 μV	–	–	–
–	86,03 μV	–	–
–	–	87,09 μV	–
–	–	–	84,69 μV
96,38 μV	91,83 μV	–	–
107,4 μV	–	108,4 μV	–
105,75 μV	–	–	105,75 μV
–	93,9 μV	108,75 μV	–
–	93,53 μV	–	95,87 μV
–	–	106,95 μV	100,25 μV
123,03 μV	107,63 μV	126,71 μV	–
121,96 μV	113,07 μV	–	101,91 μV
–	119,79 μV	125 μV	128,52 μV
123,9 μV	114,54 μV	129,09 μV	126,21 μV

zdaniem jest świetnym posunięciem. Rigol boryka się też z bardziej irytującymi wadami oprogramowania w DS1054Z, jak dziwne ucinanie śladów przy zmianie skali pionowej. Z kolei Hantek, Owon czy Uni-T nie mają niczego porównywalnego w tym segmencie cenowym, a ich tańsze oscyloskopy są ograniczone do dwóch kanałów, choć na przykład Hantek DSO2D15 w zamian oferuje szersze pasmo i wbudowany generator funkcyjny DDS w cenie około 700 zł niższej. Jedyną konkurencją dla Siglenta mogą być tylko oscyloskopy na USB, ale one borykają się ze swoimi problemami i ograniczoną funkcjonalnością.

Siglent SDS1104X-U jest w tej chwili najlepszym oscyloskopem cyfrowym na polskim rynku, oferującym najwięcej możliwości w cenie poniżej 2100 zł i niewymagającym przy tym żadnych kombinacji ze strony użytkownika. Wszystkie jego wady są zasadniczo kosmetyczne, a zatem mają znikomy wpływ na użytkowanie. Przez półtora roku użytkowania SDS1104X-U nie sprawiał mi żadnych problemów. To po prostu świetne narzędzie w świetnej cenie. Jeśli, drogi czytelniku, potrzebujesz oscyloskopu czterokanałowego, to Siglent SDS1104X-U jest najlepszym wyborem.

Paweł Kowalczyk
urgon@wp.pl

<http://www.ep.com.pl/EPwtoku>



GPSDO – jak GPS pilnuje czasu

W dzisiejszym dynamicznym i rozwijającym się świecie pomiar czasu stał się kluczowym elementem zarówno w dziedzinie nauki, technologii, jak i codziennego życia. Niezależnie od obszaru działania, od nauki po przemysł, systemy pomiaru czasu odgrywają fundamentalną rolę w precyzyjnym monitorowaniu, synchronizacji i koordynacji różnych procesów. W zaprezentowanym artykule opisano, w jaki sposób używa się systemów nawigacji satelitarnej do kontroli lokalnych zegarów.

Oscylatory GPSDO (*GPS Disciplined Oscillator*) to urządzenia elektroniczne używane do generowania stabilnego sygnału częstotliwości odniesienia, który jest zsynchronizowany z sygnałem systemu nawigacji satelitarnej takim jak GPS. Dzięki tej synchronizacji możliwe jest osiągnięcie bardzo wysokiej precyzji częstotliwości takiego zegara. Jest to istotne w aplikacjach metrologicznych czy synchronizacji czasu w wielu punktach na świecie.

Coraz większa liczba laboratoriów metrologicznych zaczyna stosować GPSDO jako swoje główne standardy częstotliwości. GPSDO ma przewagę nad standardami atomowymi, kosztując znacznie mniej oraz pełniąc funkcję „samokalibrującego” źródła czasu, które nie powinno wymagać regulacji ani okresowej kalibracji. Te cechy sprawiają, że są one atrakcyjnym wyborem dla wielu laboratoriów, ale również warsztatów elektronicznych czy hobbystycznych, jeśli potrzebne są tam tak dokładne wzorce. Niemniej jednak, źródła te mają też swoje wady w niektórych zastosowaniach.

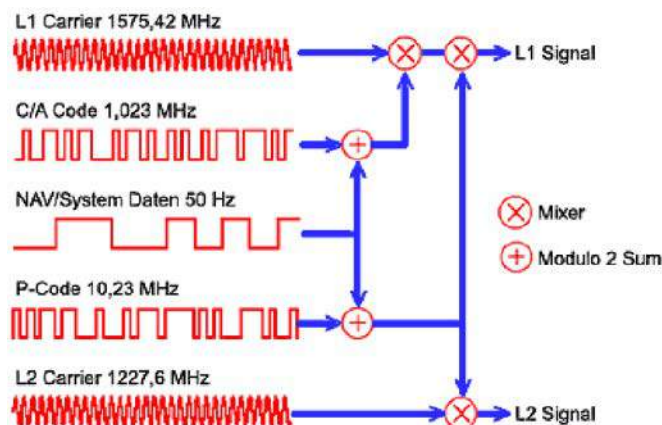
Zaprezentowany artykuł ma pełnić funkcję częściowego wprowadzenia w świat precyzyjnych i dokładnych źródeł odniesienia czasu, a równocześnie ma opisać szczegółowo jedno z nich – oscylatory „discyplinowane” za pomocą sygnału GPS. Chodzi tutaj o zegary, które są kontrolowane lub sterowane przez sygnały pochodzące z systemu

nawigacji satelitarnej po to, aby utrzymać wysoką dokładność częstotliwości. W artykule zostanie omówione, jak działają GPSDO, jak można uzyskać identyfikowalność (*traceability*) takiej kalibracji oraz zależność dokładności i precyzji tego pomiaru w zależności od modelu zegara oraz parametrów zewnętrznych.

GPS jako źródło czasu

GPS jest dobrze znanym systemem i wszechstronnym narzędziem do pozycjonowania i nawigacji. Jest również głównym systemem używanym do rozprowadzania na całym świecie czasu i częstotliwości o wysokiej dokładności. Satelity GPS są kontrolowane i obsługiwane przez Departament Obrony Stanów Zjednoczonych. Konstelacja GPS zawsze obejmuje co najmniej 24 satelity. Okrążają one Ziemię na wysokości 20 200 km w sześciu stałych płaszczyznach nachylonych pod kątem 55° do równika. Okres orbitalny wynosi 11 godzin i 58 minut (połowa długości doby gwiazdowej), co oznacza, że każdy satelita przechodzi nad danym miejscem na Ziemi o cztery minuty wcześniej niż poprzedniego dnia. Poprzez przetwarzanie sygnałów odbieranych z satelitów nawet niedrogi ręczny odbiornik GPS może określić swoją pozycję z niepewnością rzędu zaledwie kilku metrów.

Satelity GPS mają na pokładzie zegary atomowe – typowo cztery sztuki – dwa zegary cezowe i dwa zegary rubidowe; satelity konstelacji Galileo mają z kolei dwa zegary rubidowe i dwa bazujące na maserze wodorowym. Zegary te zapewniają krótkoczasowo bardzo wysoką stabilność, jednak muszą być okresowo synchronizowane z naziemnymi precyzyjnymi zegarami odniesienia. W przypadku GPS odpowiedzialne są za to naziemne stacje kontrolne Departamentu Obrony Stanów Zjednoczonych, które gwarantują, by czas satelitów zgadzał się z czasem UTC (USNO) – skalą czasu Universal Coordinated Time (UTC) utrzymywaną przez Obserwatorium Morskie Stanów Zjednoczonych (USNO). UTC (USNO) oraz czas UTC (NIST) Narodowego Instytutu Standaryzacji i Technologii (NIST) są utrzymywane w bliskiej zgodności i rzadko różnią się od siebie o więcej niż 20 ns.



Rysunek 1. Struktura sygnału GPS

Średnie przesunięcie częstotliwości między UTC (USNO) a UTC (NIST) zazwyczaj wynosi kilka części na 10^{15} lub mniej w ciągu miesiąca. NIST raportuje to przesunięcie co tydzień na swojej stronie (link znajduje się na końcu artykułu w bibliografii).

Obecnie satelity GPS nadają na dwóch częstotliwościach nośnych: L1 o częstotliwości 1,57542 GHz i L2 o częstotliwości 1,2276 GHz. Przyszłe satelity GPS mogą nadawać również inne dodatkowe częstotliwości nośne. Inne konstelacje satelitów nawigacyjnych nadawać mogą na tych samych bądź innych, np. w przypadku Galileo nadaje on w dwóch pasmach podstawowych E1 i E6, które odpowiadają pasmom L1 i L2 systemu GPS; dodatkowo Galileo nadaje sygnał E5, na który składają się pasma E5 (1,191795 GHz), E5a (1,17645 GHz) oraz E5b (1,20714 GHz).

Każdy satelita nadaje rozproszony sygnał, zwany kodem pseudolosowym (PRN), na częstotliwościach L1 i L2, i każdy satelita jest identyfikowany przez kod PRN, który nadaje. Istnieją dwa rodzaje kodów PRN. Pierwszy rodzaj to kod zgrubny (C/A) o prędkości 1023 chipów na milisekundę (1,023 megabitów/s). Drugi rodzaj to kod precyzyjny (P) o prędkości 10230 chipów na milisekundę (10,230 megabitów/s). Kod C/A jest nadawany na nośnej L1, a kod P jest nadawany na obu częstotliwościach L1 i L2. Na obu nośnych nadawana jest również wiadomość z danymi o prędkości 50 bitów/s. Struktura sygnału GPS pokazana jest na **rysunku 1**.

Zasada działania

Podstawową funkcją modułu GPSDO jest odbieranie sygnałów z satelitów GPS i używanie informacji w nich zawartych do kontroli częstotliwości lokalnego oscylatora – najczęściej rezonatora kwarcowego. Sygnały satelitarne mogą być traktowane jako odniesienie z wielu powodów. Po pierwsze, pochodzą one od zegarów atomowych zainstalowanych na pokładzie satelitów, które muszą być bardzo dokładne, aby system nawigacji satelitarnej spełniał swoje zadanie jako system pozycjonowania. Aby to zilustrować, przeliczmy, jak dokładność pozycjonowania wynika z dokładności pomiaru czasu w systemie.

Satelity GPS synchronizują swoje zegary pokładowe ze źródłami czasu pochodzącymi z kontrolnych stacji naziemnych raz podczas każdego okrążenia (czyli raz na około 12 godzin). Maksymalna dopuszczalna niepewność zegarów pokładowych przełożona w niepewność pozycjonowania systemu nie może przekroczyć jednego metra. Ponieważ światło porusza się z prędkością około 3×10^8 m/s, wymaganie niepewności poniżej 1 m jest równoznaczne z niepewnością czasu poniżej 3,3 ns. Oznacza to, że aby system GPS spełniał swoje specyfikacje, zegary satelitów muszą być wystarczająco stabilne, aby utrzymać dokładność czasu z błędem mniejszym niż 3,3 ns w okresie między korektami, co przekłada się na specyfikację stabilności częstotliwościowej rzędu 6×10^{-14} .

Po drugie, jak wspomniano – zegary atomowe znajdujące się na pokładzie satelitów GPS są okresowo (co ok. 12 godz.) synchronizowane z zegarami USNO, które z kolei są synchronizowane z czasem

odniesienia NIST. Zapewnia nam to tzw. identyfikowalność naszego standardu, którego precyzja ma odniesienie do zegarów NIST.

Celem układu GPSDO jest przeniesienie inherentnej dokładności i stabilności sygnałów satelitarnych na sygnały generowane przez lokalny oscylator. Problem przekazywania czasu i częstotliwości od zdalnego zegara do oscylatora lokalnego na dużą odległość jest analizowany przez inżynierów od dziesięcioleci. Opracowanych zostało wiele sposobów rozwiązania tego problemu. Wiele z nich jest opatentowanych, a producenci komercyjnych urządzeń GPSDO rzadko ujawniają dokładnie, w jaki sposób działają ich produkty. Niemniej jednak istnieje kilka podstawowych koncepcji, które mają zastosowanie do większości projektów. W dużym uproszczeniu mówiąc, oscylator lokalny musi być kontrolowany za pomocą jednej lub więcej pętli regulacji – sprzężenia zwrotnego, z różnymi stałymi czasowymi.

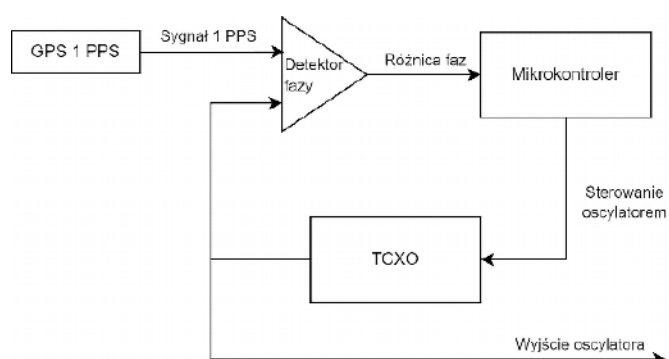
Pierwszym kluczowym elementem systemu jest odbiornik GPS. Większość producentów GPSDO korzysta z komercyjnych odbiorników GPS, ponieważ koszt opracowania takiego odbiornika jest zazwyczaj zbyt wysoki i byłoby to nieopłacalne. Oczywiście, są tutaj używane specjalne odbiorniki GPS zaprojektowane do zastosowań w zakresie synchronizacji czasu i kalibracji częstotliwości (czasami nazywane „silnikami czasu GPS”); korzystają one z wieloletnich badań i rozwoju, jednak mimo to są dość tanie – często kosztują poniżej 100 dolarów przy zakupie w większych ilościach.

Urządzenia te mogą śledzić typowo od 8 do 12 satelitów i generować sygnał jednego impulsu na sekundę (tzw. 1 PPS), który jest zsynchronizowany z sekundami UTC (USNO). Niemal wszystkie tego rodzaju moduły używają kodu C/A na nośnej L1 jako wejściowego sygnału referencyjnego do uzyskania sygnału 1 PPS.

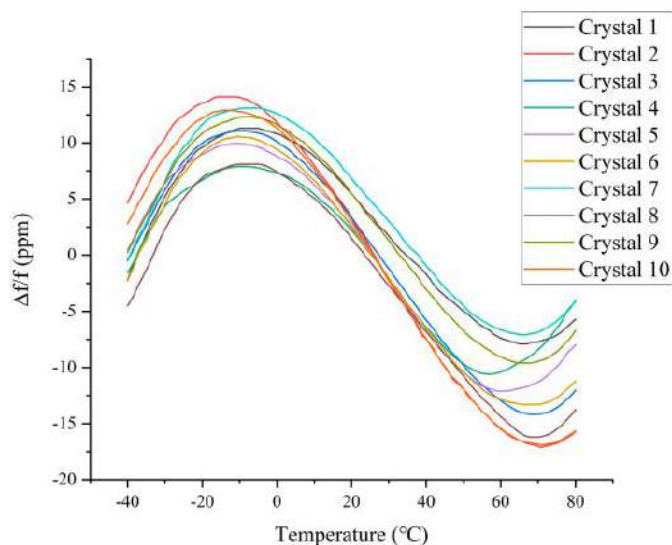
Drugim istotnym elementem układu jest oscylator lokalny. W znakomitej większości przypadków jest to zegar bazujący na oscylatorze kwarcowym. Tego rodzaju zegary są znacznie tańsze niż np. zegary atomowe – rubidowe, a do większości zastosowań laboratoryjnych sprawdzają się dostatecznie dobrze. Zwykły oscylator kwarcowy charakteryzuje się stabilnością częstotliwości na poziomie 2×10^{-5} . Dodatkowo, zegar lokalny musi mieć pewien rodzaj sterowania, który umożliwi wprowadzanie korekt do jego częstotliwości, ustalanych za pomocą sygnału GPS.

W tym celu najpraktyczniej jest zastosować TCXO, czyli oscylator kwarcowy z kompensacją temperatury, czasami nazywany oscylatorem piecowym. Jest to oscylator kwarcowy, który umieszczony jest w tzw. piecu, czyli izolowanej komorze ze stabilizowaną wewnętrzną temperaturą (w komorze znajduje się sensor temperatury oraz grzałka, a stabilizacją temperatury zajmuje się zewnętrzny układ analogowy lub mikrokontroler). W przypadku TCXO, stosowanych w tego rodzaju układach, możliwe jest zadawanie temperatury, w jakiej ma się znajdować oscylator. Typowe TCXO, bez dyscyplinowania ich z pomocą GPS, charakteryzują się stabilnością na poziomie 5×10^{-8} . Dla porównania, typowa stabilność komercyjnego zegara rubidowego wynosi 10^{-10} . Jak widać, poziom ten jest znacznie poniżej tego, co uzyskuje system GPS.

Najprostszą wersję GPSDO, której schemat blokowy pokazano na **rysunku 2**, można zbudować, stosując detektor różnicy faz do pomiaru



Rysunek 2. Schemat blokowy najprostszej implementacji GPSDO

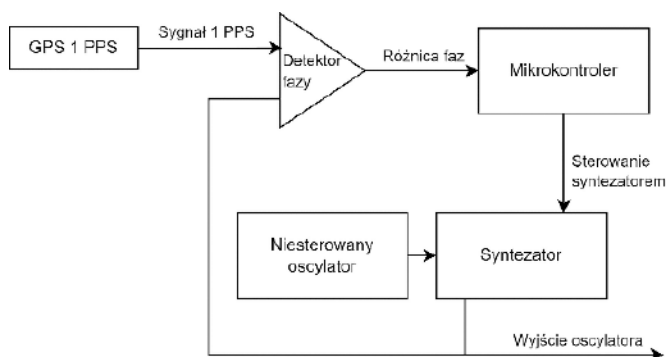


Rysunek 3. Zależność zmiany częstotliwości oscylatora kwarcowego od temperatury dla 10 przykładowych oscylatorów

różnicy między sygnałem 1 PPS z odbiornika GPS a sygnałem z zegara lokalnego. Następnie np. mikrokontroler odczytuje wyjście z detektora fazy, monitoruje różnicę faz sygnału odniesienia i lokalnego. Zadaniem układu jest utrzymanie zbliżonej do zera różnicy fazy. Kiedy się ona zmienia, np. na skutek dryftu czasowego oscylatora, oprogramowanie zmienia temperaturę TCXO, aby odstroić jego częstotliwość oscylacji. Nie jest to trywialne, ponieważ zależność zmiany częstotliwości oscylacji oscylatora kwarcowego nie jest liniowa w funkcji temperatury (**rysunek 3**). Mikrokontroler musi uwzględniać tę nieliniowość przy sterowaniu temperaturą pieca.

Producenci GPSDO stosują różne dodatkowe mechanizmy zwiększania precyzji tych systemów. Opracowano na przykład algorytm adaptacyjny, które pozwalają dodatkowo kompensować starzenie się oscylatora kwarcowego i zmiany temperatury zewnętrznej dla wielu rodzajów oscylatorów lokalnych, co pozwala na skuteczną pracę z tanimi zegarami kwarcowymi. Niektóre algorytmy nawet „uczą się” charakterystyki oscylatora – nie jest ona taka sama dla każdego elementu, podobny jest tylko w zasadzie kształt. Pozwala to na łatwiejsze sterowanie zmianami częstotliwości oscylatora lokalnego, pozwalając na lepsze sterowanie wyjściem GPSDO, w szczególności w sytuacjach, gdy sygnał wejściowy GPS zostanie tymczasowo utracony. To zapewnia GPSDO zdolność do utrzymania dokładnej częstotliwości przez pewien czas np. w przypadku odłączenia anteny lub utraty widoczności satelitów.

Inny rodzaj systemu GPSDO pokazano na **rysunku 4**. Układ ten nie koryguje częstotliwości oscylatora lokalnego. Zamiast tego wyjście swobodnie działającego oscylatora lokalnego jest wysyłane do syntezy częstotliwości. Korekty częstotliwości są następnie stosowane do wyjścia syntezy. Współczesne bezpośrednie syntezy



Rysunek 4. Schemat blokowy implementacji GPSDO z zastosowaniem syntezy DDS



Fotografia 1. Niewielka antena GPS musi być umieszczona w miejscu z widokiem na niebo, np. na dachu

cyfrowe (DDS) mają bardzo wysoką rozdzielczość i pozwalają na dokonywanie bardzo małych korekt częstotliwości. Na przykład DDS o rozdzielczości 48 bitów może zapewnić rozdzielczość częstotliwości poniżej mikroherca przy 10 MHz (taka rozdzielczość pozwala na korekty z dokładnością na poziomie ok. $3,5 \times 10^{-8}$ Hz). Ponadto użycie lokalnego oscylatora bez sterowania często przynosi lepsze wyniki niż sterowanie nim, gdyż nieoczekiwane zmiany sygnału sterującego mogą prowadzić do niepożądanych zmian jego częstotliwości wyjściowej.

Praktyczne implementacje

GPSDO to zaawansowane urządzenia, których opracowanie wymagało ogromnego wysiłku inżynierskiego. Niemniej jednak ich instalacja i użytkowanie jest stosunkowo proste dla użytkownika końcowego. Najtrudniejszą częścią instalacji jest zamocowanie anteny systemu GPS na dachu (na **fotografii 1** pokazano przykładową antenę na dachu autora) w miejscu z wolnym widokiem na niebo. Dodatkowo, antenę należy umieścić stosunkowo blisko naszego warsztatu, aby zminimalizować straty sygnału wzdłuż kabla łączącego antenę GPS z odbiornikiem w zegarze. Po zainstalowaniu GPSDO system zazwyczaj rozpoczyna badanie pozycji anteny od razu po włączeniu. Badanie to jest procesem jednorazowym, który zazwyczaj trwa kilka godzin od włączenia urządzenia. Po zakończeniu badania pozycji anteny i synchronizacji z satelitami GPSDO urządzenie jest gotowe do użycia jako standard częstotliwości i czasu.

Większość modułów GPSDO generuje na wyjściu sygnały sinusoidalne o częstotliwości 10 MHz jako sygnał odniesienia częstotliwości, a także ma wyjście sygnału 1 PPS jako odniesienie interwałów czasowych oraz do synchronizacji czasu z UTC.



Fotografia 2. Tylny panel urządzenia GPSDO

Fotografia 2 pokazuje zdjęcie tylnego panelu przykładowego urządzenia GPSDO. Ma ono dodatkowe dwa wyjścia cyfrowe z zegarem 10 MHz w standardzie High Speed CMOS (HCMOS), które jest przeznaczone do taktowania układów cyfrowych. Oprócz tych wyjść na tylnym panelu znajduje się miejsce do podłączenia zdalnej, aktywnej anteny GPS, port RS232, zapewniający możliwość konfiguracji czy bardziej zaawansowanej kontroli modułu oraz, oczywiście, port zasilania.

W przypadku tego konkretnego modelu znajduje się w nim prosty wbudowany wzmacniacz dystrybucyjny dla sygnałów HCMOS. Jednak, jeśli w naszej instalacji w laboratorium potrzebne są więcej niż 2 cyfrowe sygnały zegarowe odniesienia, to konieczne jest dodanie do systemu zewnętrznego wzmacniacza dystrybucyjnego. Jest to prosty, aktywny rozdzielacz sygnału, który charakteryzuje się niskim szumem fazowym, co zapewnia zachowanie wysokiej dokładności GPSDO.

Podsumowanie

Powyższy artykuł poświęcony został opisowi zasady działania i zastosowaniom układów GPSDO. Są one niezwykle istotne dla inżynierów elektroników, którzy zajmują się testowaniem i uruchamianiem urządzeń elektronicznych. GPSDO to zaawansowane urządzenia, które korzystają z sygnałów z systemu GPS do utrzymania bardzo stabilnej i dokładnej częstotliwości oscylatora lokalnego. Dzięki zastosowaniu relatywnie niedrogich oscylatorów kwarcowych sterowanych termicznie (TCXO) lub scalonych, cyfrowych syntezyatorów (DDS) urządzenia GPSDO są niedrogie w porównaniu do innych, precyzyjnych źródeł zegara odniesienia.

Precyzyjny pomiar czasu jest kluczowy w laboratoriach elektronicznych, a układy GPSDO dostarczają możliwości, które są niezbędne w wielu aplikacjach. Po pierwsze, GPSDO umożliwiają dokładne i stabilne generowanie częstotliwości, co jest kluczowe w zastosowaniach

związanych z wieloma pomiarami elektronicznymi. Możliwość uzyskania stabilnych sygnałów o znanej częstotliwości jest nieoceniona w takich dziedzinach jak metrologia, badania mikrofalowe czy też np. badania materiałów.

Po drugie, układy GPSDO są niezastąpione w systemach synchronizacji czasu. Dzięki nim można utrzymywać dokładną i zsynchronizowaną w czasie pracę wielu urządzeń rozproszonych nawet na dużej odległości, co jest kluczowe w telekomunikacji, sieciach komputerowych czy w niektórych badaniach naukowych (np. z tego rodzaju synchronizacji korzystają akceleratory cząstek elementarnych).

Warto podkreślić, że układy GPSDO są jednymi z najtańszych narzędzi w laboratoriach elektronicznych, które pozwalają na precyzyjny pomiar czasu czy częstotliwości. Dają one możliwość uzyskania wyników pomiarów na najwyższym poziomie dokładności i precyzji, co jest kluczowe w wielu dziedzinach inżynierii, przy tym kosztując ułamek ceny np. zegarów atomowych.

Nikodem Czechowski, EP

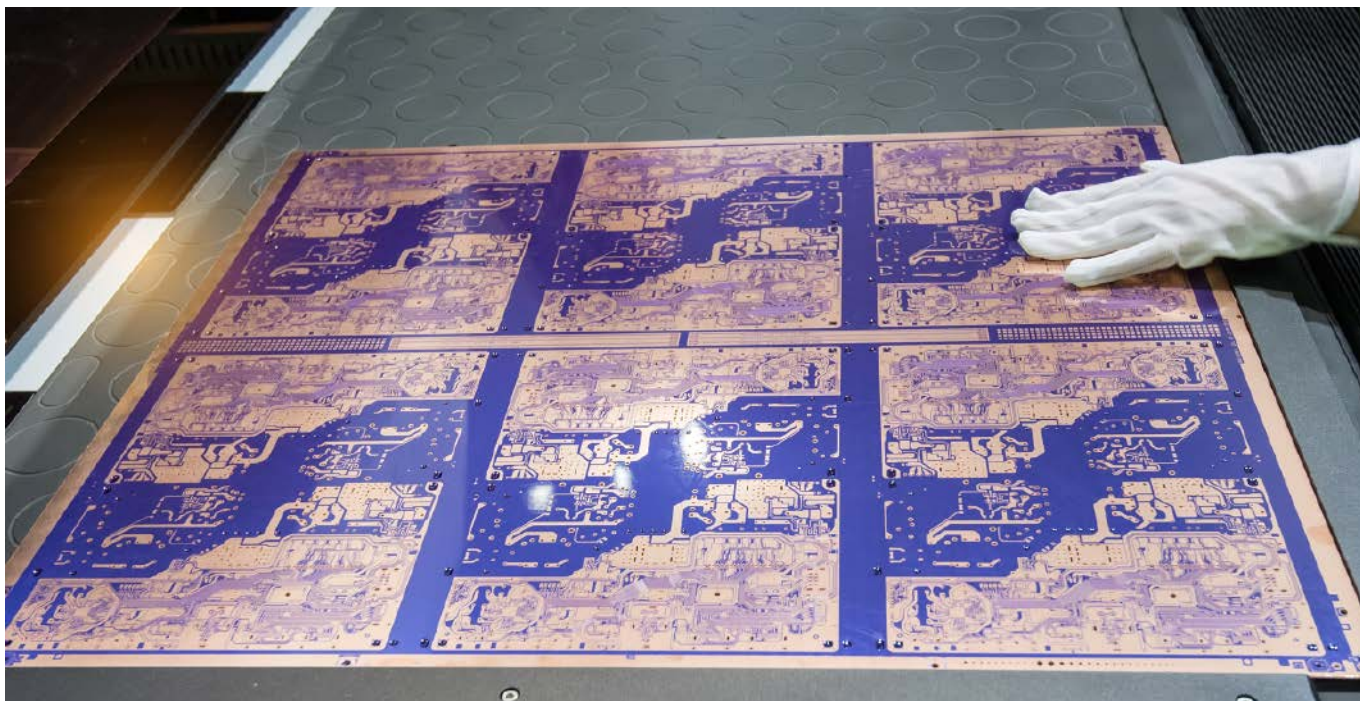
Bibliografia

1. D. Shah „GPS disciplined oscillator”, Biley Technologies whitepaper, 2017.
2. H. El-Bakry i N. Mastorakis „Design of anti-GPS for reasons of security”, Proceedings of the International Conference on Computational and Information Science 2009.
3. M. Lombardi „The Use of GPS Disciplined Oscillators as Primary Frequency Standards for Calibration and Metrology Laboratories”, NCSLI Measure: The Journal of Measurement Science 3 (2008).
4. <http://tf.nist.gov/pubs/bulletin/nistusno.htm>
5. Z. Wang, J. Wu „A Method to Increase the Frequency Stability of a TCXO by Compensating Thermal Hysteresis”, Sensors 20 (2020)

REKLAMA

Elektronika od podstaw do praktyki	Tematy	Posty	Ostatni post	Ostatnie posty
1. Elektronika - tematy dowolne Tematy ogólne związane z elektroniką. Dyskusja nt. podzespołów, zasad działania komponentów itp. Moderatorzy: Jacek Bogusz, Moderatorzy	5109	26678	Re: Okap czy pochłaniacz autor: cefk G 20 lis 2020, o 08:44	wczoraj, o 16:30 Czesiek: Mój zdaniem warto po prostu rozstać z anajomym, który rozstał swoim znajomym, którzy... Wiesz o co mi chodzi, to naprawdę szybko działa. Dobr
2. Serwis urządzeń elektronicznych Pytania i porady dotyczące serwisu urządzeń elektronicznych Moderatorzy: Jacek Bogusz, Grzegorz Becker, Moderatorzy	1121	4799	dymomierz autor: miki-kajak G 14 lut 2020, o 13:02	wczoraj, o 07:30 zidane: Drukarka Fingerprint. Współpracuję z nimi od dłuższego czasu. Dobra drukarka z którą w firmie współpracuję już od dawna. W ofercie ma a i z ad
3. Aparatura kontrolno-pomiarowa i narzędzie Wszystko na temat aparatury kontrolno-pomiarowej oraz	31	179	Gdzie dostępne bezpieczniki 10kV autor: porlock G	

O projektach, miniprojektach, projektach soft i na wiele innych tematów dyskutuj na forum.ep.com.pl



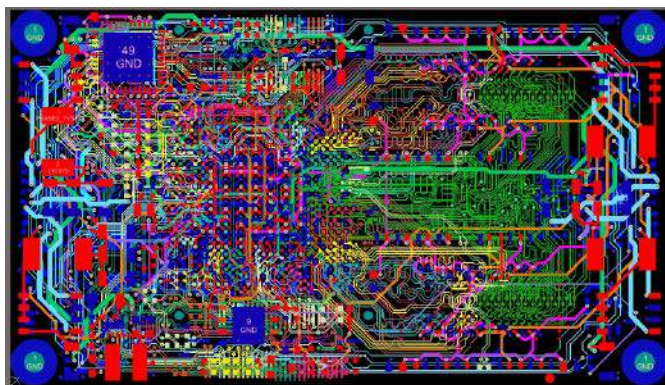
Zagadnienia materiałowe w produkcji wielowarstwowych obwodów drukowanych

Współczesna elektronika, w swoim „odwiecznym” pędzie do miniaturyzacji, nie mogłaby istnieć bez wielowarstwowych obwodów drukowanych, pozwalających na układanie elementów przy pozostawieniu jedynie minimalnych odstępów pomiędzy padami sąsiadujących podzespołów. Gęstość upakowania komponentów i wynikające z niej możliwości miniaturyzacji to zaledwie jedna strona medalu – właściwe żonglowanie parametrami technologicznymi stosu warstw oraz prawidłowe użycie technik projektowych mają także ogromny wpływ na osiągi urządzenia w zakresie emisji i podatności elektromagnetycznej, stabilności napięcia zasilania, integralności sygnałów w szybkich torach analogowych i cyfrowych czy wreszcie... parametrów termicznych. Dokonamy przeglądu najważniejszych zagadnień związanych tak z samym procesem produkcji płytek wielowarstwowych, jak i prawidłowym przygotowaniem projektu. Aby jednak w pełni zrozumieć technologię MLPCB, warto najpierw zapoznać się bliżej z aspektami materiałowymi.

„Plusy dodatnie i plusy ujemne”, czyli zalety i wady obwodów wielowarstwowych

Główną motywacją do opracowania i upowszechnienia technologii obwodów wielowarstwowych (Multi-Layer Printed Circuit Board, MLPCB) była rosnąca liczba komponentów elektronicznych, montowanych w dodatku na coraz mniejszych płytkach drukowanych. Nie da się ukryć, że w największym stopniu przysłużyły się do tego urządzenia cyfrowe, w tym przede wszystkim płyty główne komputerów (**rysunek 1**), karty graficzne i inne rozbudowane podsystemy, których nie sposób było produkować z użyciem konwencjonalnych obwodów dwuwarstwowych. Dziś wielowarstwowe PCB są już niekwestionowanym standardem – każdy ma w swojej kieszeni smartfon, którego płyta główna ma zwykle 10, a nawet 12 warstw, a 4-warstwowe PCB coraz częściej wypierają dwuwarstwowe w konstrukcjach o mniejszej gęstości upakowania elementów.

Długą listę zalet MLPCB otwiera rzecz jasna możliwość miniaturyzacji urządzeń elektronicznych, a zaraz za nią należy wymienić łatwość prowadzenia dużej liczby połączeń w różnych kierunkach (bez konieczności nadmiernego lawirowania w celu uniknięcia przecięć z już poprowadzonymi ścieżkami) oraz doskonałe osiągi w zakresie ekranowania wrażliwych linii sygnałowych przez wewnętrzne płaszczyzny masy. Zastosowanie równoległych płaszczyzn połączonych z głównymi szynami zasilania pozwala ponadto na uzyskanie



Rysunek 1. Przykładowy projekt wielowarstwowej płyty głównej minikomputera jedno płytowego (<https://bit.ly/3tNgYhF>)

doskonałych warunków pracy układu, co wiąże się zarówno z niską impedancją połączeń, jak i obecnością pojemności rozproszonej, którą wykazują (przeźrzone zaledwie cienką warstwą dielektryka) pola masy i zasilania. Mało tego – płyty wielowarstwowe wykazują większą sztywność mechaniczną w porównaniu do jedno- i dwuwarstwowych, a ich parametry termiczne mogą być kształtowane w znacznie szerszym stopniu, niż miałyby to miejsce w przypadku PCB o prostszej konstrukcji stosu. Zaawansowane techniki projektowania obwodów wielowarstwowych dają ponadto szansę na implementację struktur funkcjonalnych, których wykonanie na płytkach jedno- lub dwuwarstwowych byłoby albo niemożliwe, albo całkowicie nieefektywne.

Mniej oczywistą, choć także mającą praktyczne znaczenie zaletą płyt wielowarstwowych, jest znaczące utrudnienie dla osób próbujących zrekonstruować układ połączeń na PCB za pomocą inżynierii odwrotnej. Każdy, kto dokonywał takiej operacji na płycie dwuwarstwowej, doskonale zna metody i narzędzia, które mogą być zastosowane do odwzorowania schematu lub jego części. W przypadku płyt wielowarstwowych sprawa okazuje się wyraźnie trudniejsza, wymusza bowiem... mechaniczne usuwanie kolejnych warstw miedzi, prepregów oraz rdzeni (fotografia 1), co nie jest już zadaniem trywialnym. I choć nie ma na świecie urządzenia, którego nie dałoby się skopiować (o czym świadczą chociażby liczne klony markowych sprzętów codziennego użytku), to zastosowanie druku wielowarstwowego może choć częściowo utrudnić i wydłużyć cały proces, zniechęcając przynajmniej tych słabiej zdeterminowanych „zawodników”.

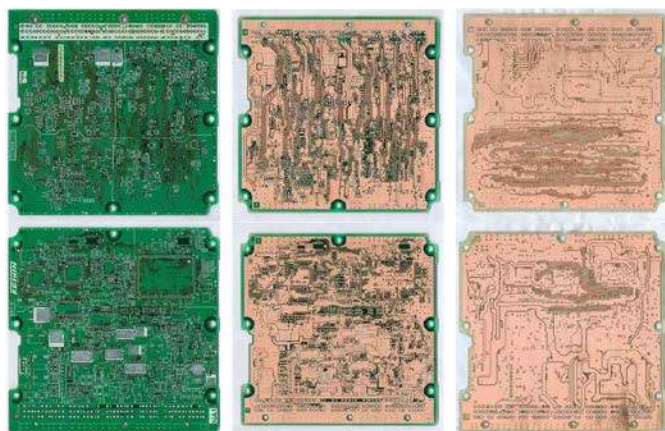
Decydując się na zastosowanie druku wielowarstwowego, nie należy jednak zapominać o pewnych ograniczeniach – oprócz znacznie wyższej ceny produkcji, dłuższego czasu realizacji zamówień oraz wyższego stopnia złożoności technologii MLPCB (który wciąż wyłącza z rynku niektórych producentów PCB o mniej rozbudowanym parku maszynowym), należy brać także pod uwagę trudności



Fotografia 2. Przykład naprawy płyty wielowarstwowej – tymczasowe usunięcie ścieżki w wyżej położonej warstwie w celu uzyskania dostępu do warstwy głębszej (<https://bit.ly/46EiSe2>)

z diagnozowaniem usterek oraz dokonywaniem inspekcji. Podczas gdy awarie płyt dwuwarstwowych udaje się naprawić niemal zawsze (choć – w zależności od posiadanego doświadczenia, umiejętności i sprzętu – z bardzo zróżnicowanymi efektami końcowymi), to już w przypadku obwodów wielowarstwowych problem staje się znacznie poważniejszy. W praktyce, wszelkie modyfikacje oraz pomiary mogą być prowadzone wyłącznie na warstwach zewnętrznych, choć w sieci można znaleźć (nieco wiekowe) materiały, opisujące karkołomne techniki ratowania uszkodzonych płyt wielowarstwowych za pomocą technik „odkrywkowych” z użyciem precyzyjnej minifrezarki, skalpela chirurgicznego i specjalnych, prefabrykowanych elementów naprawczych (fotografia 2). Zwiększona pojemność cieplna może ponadto utrudniać zarówno montaż rozplądowy, jak i lutowanie punktowe oraz wylutowywanie elementów.

Warto zwrócić uwagę na fakt, że w codziennej praktyce często nie zastanawiamy się głębiej nad mikrostrukturą materiałów, używanych do produkcji PCB, a także nad ich licznymi parametrami, spotykanymi w katalogach producentów. Mało tego – nawet w wielu materiałach branżowych możemy spotkać się z ogólnikami w postaci stwierdzeń, że „parametry płytek zależą od rodzaju laminatu”, a w opozycji do klasycznego FR-4 stawia się m.in. materiały marki Rogers, dodając jedynie zdawkową informację, że oferują one lepsze parametry w obwodach RF. Tymczasem każdy z producentów nowoczesnych laminatów i prepregów dokłada wszelkich starań, by sukcesywnie poszerzać swoją ofertę o coraz bardziej zaawansowane rodzaje materiałów, dostosowane do pracy w najróżniejszych warunkach środowiska i rodzajach obwodów. Dlatego też w dalszej części artykułu przyjrzymy się często pomijanym aspektom – wyjaśnimy m.in., co oznaczają poszczególne oznaczenia prepregów, dokładnie wyjaśnimy budowę spłotów włókna szklanego, opiszemy także najważniejsze parametry laminatów oraz defekty materiałowe, będące przyczyną awarii płytek wielowarstwowych.



Fotografia 1. Czterowarstwowa płyta drukowana oraz widoki poszczególnych warstw, odstawianych sukcesywnie metodą szlifowania w ramach procesu inżynierii odwrotnej (<https://bit.ly/3tMrT6g>)

REKLAMA

LASEROWE SZABLONY DO MONTAŻU SMT

Materiał: stal nierdzewna CrNi
Zakres grubości blach: 0,020–1,000 mm
Wycinamy również detale
o dowolnych kształtach



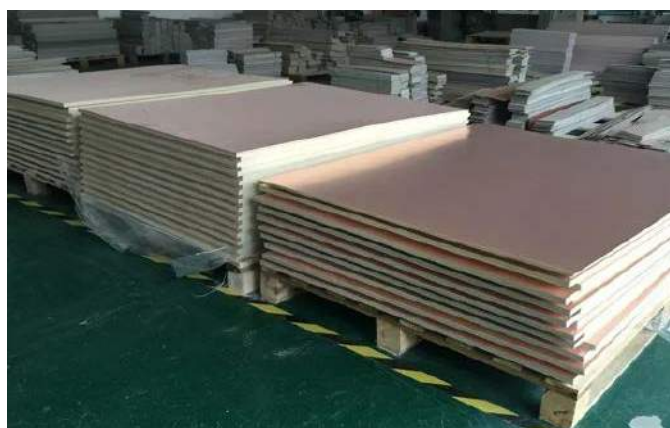
LASTENIC LASER & ELECTRONICS sp. z o.o.
58-100 Świdnica, ul. Husarska 5
tel. 74 851 48 77, 697 977 732
www.lastenic.com info@lastenic.com

Layer	Stack up	Description	Material	Thickness	Type	Supplier	Quantity	Unit
1	Liquid PhotoImageable Mask		SolderMask	18.000	Prepreg	PERI		
2	Copper Foil		Copper	35.000	Copper	PERI		
3	FR4 Core		FR4	76.140	Core	PERI		
4	Prepreg 7628		Prepreg	35.000	Prepreg	PERI		
5	Copper Foil		Copper	35.000	Copper	PERI		
6	FR4 Core		FR4	76.140	Core	PERI		
7	Prepreg 7628		Prepreg	35.000	Prepreg	PERI		
8	Copper Foil		Copper	35.000	Copper	PERI		
9	Liquid PhotoImageable Mask		SolderMask	18.000	Prepreg	PERI		

Rysunek 2. Przykładowy projekt stosu 12-warstwowej PCB o grubości wynikowej 1,6 mm (<https://bit.ly/3seXr4a>)



Fotografia 3. Proces ręcznego składania stosu PCB z poszczególnych rdzeni i prepregów (<https://bit.ly/3QcqsFD>)



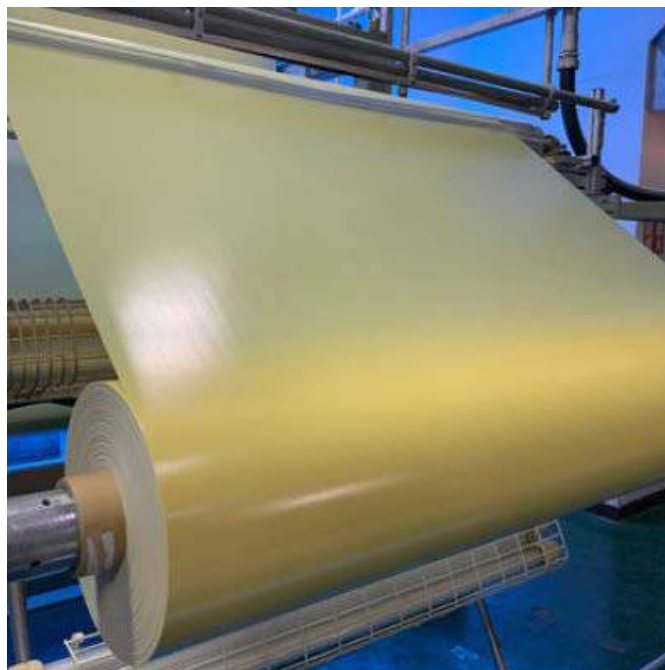
Fotografia 4. Wnętrze przemysłowego magazynu laminatów (<https://bit.ly/3ShqTRN>)

Budowa stosu – podstawy technologiczne

Kluczowym zagadnieniem w projektowaniu płytek wielowarstwowych jest zastosowana przez konstruktora budowa stosu. To właśnie ona – obok rodzaju stosowanych materiałów – wpływa bowiem w największym stopniu na niemal wszystkie parametry technologiczne obwodu drukowanego.

Jedną z podstawowych zasad dobrej praktyki w projektowaniu MLPCB jest zachowanie możliwie symetrycznej konstrukcji stosu, z parzystą liczbą warstw miedzi. Choć wykonanie PCB o 3, 5 czy np. 9 warstwach jest możliwe, to stanowi ono pewne wyzwanie technologiczne i – o ile producent podejmie się takiego zdania – realizacja projektu będzie wyraźnie droższa niż w przypadku konwencjonalnej płytki o nawet większej, ale parzystej liczbie warstw. Nie warto więc brać za przykład płytki, która – wyprodukowana w Japonii w 2012 roku – pobila rekord Guinnessa dzięki największej w historii PCB liczbie warstw, równej – bagietela – 129 (!).

W produkcji płyt wielowarstwowych stosowane są trzy podstawowe rodzaje prefabrykatów – rdzeń (*core*), prepreg oraz folia miedziana (*copper foil*) – patrz **rysunek 2** i **fotografia 3**. Rdzeń pełni funkcję bazy



Fotografia 5. Produkcja prepregu (<https://bit.ly/3MfdMg1>)

Layer	Stack up	Description	Type
1	Liquid PhotoImageable Mask	SolderMask	
2	Copper Foil	Copper	
3	Prepreg 7628	Dielectric	
4	FR4 Core	FR4	
5	Prepreg 7628	Dielectric	
6	Copper Foil	Copper	
7	Liquid PhotoImageable Mask	SolderMask	

Rysunek 3. Stos typu *foil build* (<https://bit.ly/45Kr23s>)

konstrukcyjnej, na której opiera się cała płyta lub jej część – w większości typowych aplikacji ma on postać klasycznego, cienkiego laminatu szklano-epoksydowego, pokrytego obustronnie folią miedzianą (**fotografia 4**). Rzecz jasna, możliwe jest także wykonanie rdzeni z wielu innych rodzajów materiału, jednak zawsze rola tegoż elementu pozostaje taka sama – stanowi on nośnik warstw sygnałowych lub płaszczyzn masy/zasilania, a jednocześnie bierze istotny udział w nadawaniu płytce odpowiedniej sztywności (choć cienkie rdzenie przypominają bardziej folię niż sztywny laminat, znany ze zwykłych płytek 1- i 2-warstwowych o grubości 1,0 mm czy też 1,6 mm). Prepreg (**fotografia 5**) to najczęściej także mata z włókna szklanego, nasączona żywicą epoksydową z domieszką różnego rodzaju wypełniaczy (których celem jest m.in. poprawa przewodności cieplnej, zwiększenie wytrzymałości na zginanie czy też osiągnięcie odpowiedniej klasy palności) – w tym przypadku jednak epoksyd nie jest fabrycznie utwardzony, dzięki czemu w procesie prasowania może on trwale związać się z powierzchniami sąsiadujących rdzeni, prepregów bądź folii miedzianej. Nawiasem mówiąc, składanie dwóch lub trzech czy nawet czterech prepregów, nieprzezielonych żadnym innym materiałem, stanowi klasyczną praktykę, stosowaną przez producentów w celu uzyskania odpowiedniej grubości dielektryka pomiędzy poszczególnymi warstwami miedzi. Folia miedziana znajduje natomiast zastosowanie jako zewnętrzna „okładzina” całego stosu – to właśnie ona, wytrawiona jako ostatnia, będzie tworzyła warstwy zewnętrzne (*top*, *bottom*).

Można wyróżnić dwie główne odmiany konstrukcji stosu MLPCB, określane jako *foil build* oraz *core build*, zaś rozróżnienie – jak sama nazwa wskazuje – opiera się na ułożeniu zastosowanych prefabrykatów. Na **rysunku 3** pokazano (zdecydowanie najpopularniejszą) konstrukcję typu *foil build* na przykładzie płytki 4-warstwowej. Jak widać, centralnym elementem płytki jest rdzeń, stanowiący nośnik warstw wewnętrznych, zaś po obydwu jego stronach znajdują się

Layer	Stack up	Description	Type
1		Liquid Photoimageable Mask	SolderMask
2		FR4 Core	FR4
3		PrePreg 1651	Dielectric
4		FR4 Core	FR4
		Liquid Photoimageable Mask	SolderMask

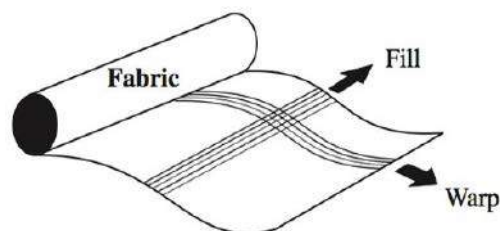
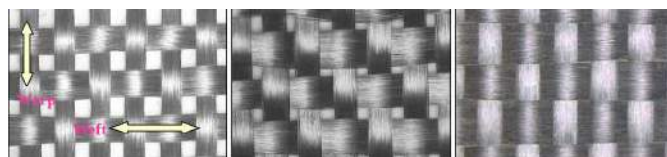
Rysunek 4. Stos typu core build (<https://bit.ly/45Kr23s>)

zestawy prepregów, „zakończone” folią miedzianą. Diametralnie inną filozofię stosują producenci, oferujący ponadto technologię *core build* – w tym przypadku „goła” folia miedziana w ogóle nie bierze udziału w procesie produkcyjnym, gdyż zewnętrzne warstwy obwodu drukowanego są w istocie częścią rdzeni (rysunek 4). Ten rodzaj stosu jest stosowany m.in. w układach mikrofalowych, bazujących na kosztownych laminatach specjalnego przeznaczenia – pary współpracujących ze sobą warstw są zatem przedzielone tym cennym dielektrykiem, podczas gdy tak wykonane płytki dwustronne łączy się następnie ze sobą za pomocą znacznie tańszych prepregów.

Prepregi i rdzenie – wgląd w parametry materiałowe i procesowe

Jak już wspomnieliśmy, najbardziej rozpowszechnioną grupą materiałów, stosowanych w produkcji płyt wielowarstwowych, są rozmaite odmiany klasycznego FR-4, którego konstrukcja bazuje na tkaninie z włókna szklanego, nasączonej żywicą epoksydową. Nic więc dziwnego, że ten sam materiał podstawowy występuje zarówno w formie w pełni utwardzonej (jako baza rdzenia, pokrywana następnie miedzią), jak i nieutwardzonej (jako gotowy do użytku prepreg). Z uwagi na inne zadania rdzeni i prepregów, charakteryzujące je parametry pokrywają się jednak ze sobą tylko w pewnej części.

Warto wiedzieć, że mechaniczne, termiczne i elektryczne właściwości prepregu zależą w dużej mierze od zastosowanego splotu włókna – i to właśnie od niego pochodzą 3- lub 4-cyfrowe oznaczenia,

Rysunek 5. Nazewnictwo nici splotu włókna w zależności od kierunków, określonych w procesie tkania (<https://bit.ly/46PWGxP>)Fotografia 6. Porównanie splotów tkaniny szklanej typu standardowego – po lewej, rozszerzonego w jednym kierunku (w środku) i mechanicznie rozszerzonego w obu kierunkach – po prawej (<https://bit.ly/3MDOQiP>)

stosowane m.in. w opisach stosu warstw płyt MLPCB. Każdy rodzaj splotu ma swój unikalny numer i ściśle określone parametry geometryczne – aby to lepiej zrozumieć, musimy jednak zajrzeć w głąb włókna szklanego za pomocą mikroskopu o 50-krotnym powiększeniu. Dość powiedzieć, że w specyfikacji poszczególnych splotów bierze się pod uwagę raster (*pitch*), wymiary przekroju „nitki” oraz liczbę zakończeń/cal (i to osobno dla obydwu kierunków, określanych jako wątek (*fill*, *weft*) i osnowa (*warp*) – patrz rysunek 5), a także określa się rodzaj splotu jako standardowy lub rozszerzony w jednym (*expanded*) bądź dwóch (*mechanical spread*) kierunkach (fotografia 6). Najważniejsze parametry popularnych splotów zestawiono w tabeli 1, zaś fenomenalnie opracowany przekrój płyty wielowarstwowej, na którym doskonale widoczne są poszczególne włókna szklane, zawieszane w półprzezroczystej żywicy, można zobaczyć na fotografii 7.

Tabela 1. Parametry „konstrukcyjne” najpopularniejszych splotów włókna szklanego (<https://bit.ly/3tLboac>)

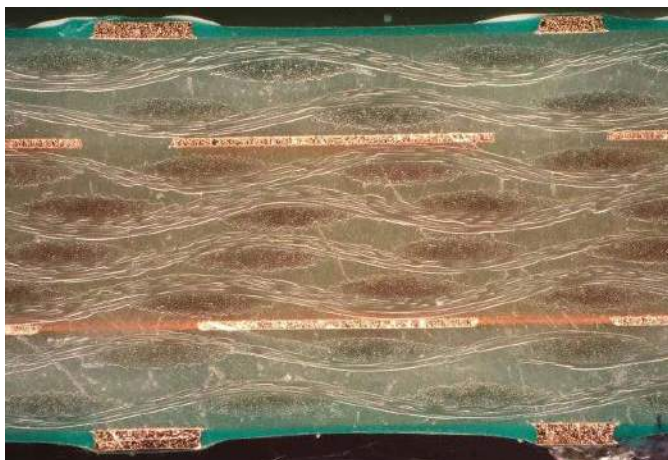
Typ splotu	Liczba zakończeń/cal w osnowie	Liczba zakończeń/cal w wątku	Grubość włókna	Grubość laminatu	Zawartość żywicy
106	56	56	0,0015	0,002	69,00%
1080	60	47	0,0025	0,003	62,00%
2113	60	56	0,0029	0,004	54,50%
3313	61	62	0,0031	0,004	54,00%
3070	70	70	0,0034	0,004	49,50%
2116	60	58	0,0038	0,005	51,80%
1652	52	52	0,0045	0,005	42,00%
7628	44	32	0,0065	0,007	44,40%

Tabela 2. Parametry płynięcia dla najpopularniejszych typów prepregów (<https://bit.ly/3s61IXE>)

Oznaczenie prepregu	Grubość nominalna		Średnia zawartość żywicy	Płynięcie (wartość skalowana, wyznaczona z pomocą prasy hydraulicznej)		Obliczona grubość (po autoklawowaniu)
	[mm]	[cal]		mil/warstwę	mm/warstwę	
104 AT 01	0,040	0,0015	72	1,5±0,2	0,038±0,005	0,041
106 AT 01	0,050	0,0020	73	1,9±0,2	0,048±0,005	0,059
1080 AT 01	0,063	0,0025	62	2,7±0,3	0,069±0,008	0,078
2113 AT 01	0,095	0,0037	54	3,7±0,3	0,094±0,008	0,101
2125 AT 01	0,100	0,0039	53	3,9±0,3	0,099±0,008	0,106
2116 AT 01	0,115	0,0045	50	4,7±0,3	0,100±0,008	0,120
2165 AT 01	0,150	0,0059	55	5,1±0,3	0,130±0,008	0,158
7628 AT 01	0,180	0,0071	45	7,0±0,3	0,178±0,008	0,201
1080 AT 02	0,063	0,0025	60	2,5±0,3	0,064±0,008	0,071

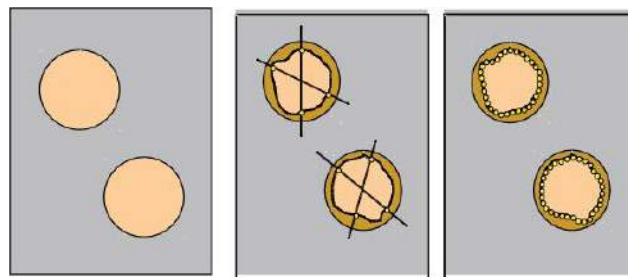
Tabela 3. Grubości prepregów w zależności od warstwy, z którą są łączone (podano wartości dla warunków, w których prepregi zawsze występują w zestawach po 2 arkusze lub więcej) – <https://bit.ly/3tRKgHa>

		Grubość warstwy prepregu [mil]							
		po połączeniu z rdzeniem pokrytym miedzią o grubości d [μm], przy pokryciu p [%]						po połączeniu z folią Cu na warstwie zewnętrznej	dla dodatkowych arkuszy prepregu niesąsiadujących z miedzią
	d	18		35		70			
	p	30	70	30	70	30	70		
Rodzaj prepregu	106	1,9	2,2	1,5	2,0	0,5	1,5	2,3	2,1
	1080	2,6	2,8	2,1	2,6	1,1	2,2	3,0	2,7
	2113	3,5	3,7	3,0	3,5	2,0	3,1	3,9	3,5
	2116	4,7	4,9	4,2	4,7	3,2	4,3	5,1	4,6
	7628	6,5	6,8	6,0	6,5	5,0	6,1	6,9	6,2

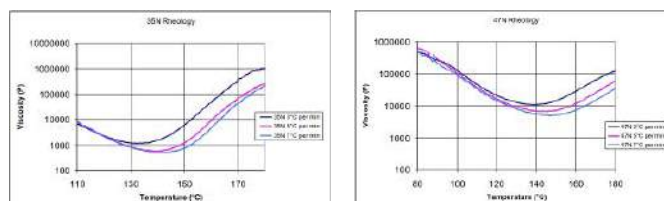


Fotografia 7. Przekrój czterowarstwowej płyty na bazie laminatu szklano-epoksydowego FR4, opracowany w laboratorium firmy Eurocircuits (<https://t.ly/CnfAr>)

Te wszystkie parametry, choć pozornie mało istotne dla konstruktora PCB, w rzeczywistości mają bardzo duży wpływ na zachowanie prepregu podczas procesu laminacji – ponieważ każdy ze spłotów ma inną grubość nominalną oraz zostawia więcej lub mniej przestrzeni, w której może pomieścić żywicę epoksydową, w efekcie poszczególne rodzaje prepregów różnią się parametrami płynięcia żywicy (*resin flow*) oraz (wyrażonej procentowo w odniesieniu do masy półproduktu) zawartości tegoż materiału (*resin content*) – patrz **tabela 2**. Należy bowiem pamiętać, że w czasie laminacji (rozgrzewania przy jednoczesnym ściskaniu stosu) część żywicy bierze udział w wytworzeniu trwałego wiązania pomiędzy prepregiem a okalającymi go warstwami. Najlepszym dowodem na „ucieczkę” pewnej części żywicy z wnętrza prepregu jest fakt, iż końcowa grubość PCB zależy



Rysunek 6. Metody pomiaru średnicy otworów podczas testu płynięcia żywicy (opis w tekście) (<https://bit.ly/476fQPK>)



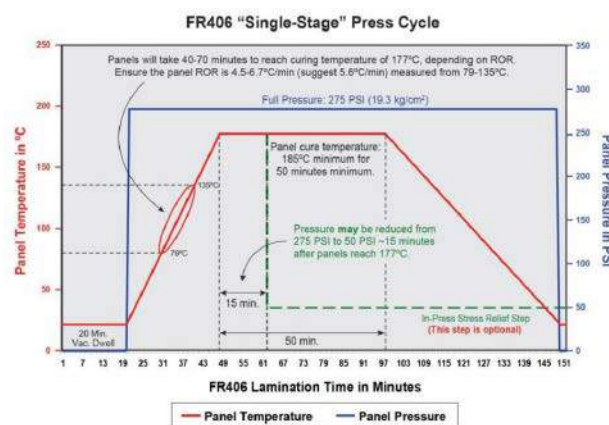
Rysunek 7. Porównanie lepkości standardowego prepregu polimernego typu 35N (na górze) z lepkością epoksydowego prepregu typu low-flow (47N) – <https://bit.ly/3s665C2>

w pewnym (choć stosunkowo niewielkim) stopniu od... procentowego pokrycia danej warstwy miedzią oraz wyjściowej grubości folii miedzianej. W przypadku płaszczyzn masy/zasilania o pokryciu rzędu 70% lub więcej wynikowa grubość połączenia będzie nieco większa niż w przypadku warstw sygnałowych o pokryciu na poziomie np. 30% – przybliżone wartości można znaleźć w **tabeli 3**.

Płynięcie żywicy jest na tyle istotnym parametrem z punktu widzenia produkcji zaawansowanych obwodów drukowanych (m.in. płyt sztywno-giętkich czy też zawierających wstawki chłodzące, tzw. *copper coins*), że parametr ten doczekał się międzynarodowych standardów (IPC TM-650 2.3.17, IPC TM-650 2.3.17.2), określających



Fotografia 8. Arkusz prepregu, przeznaczony do pomiarów płynięcia żywicy (<https://bit.ly/476fQPK>)



Rysunek 8. Profil laminacji jednostopniowej dla laminatów FR406 marki ISOLA (<https://t.ly/m6EaQ>)

Tabela 4. Parametry folii miedzianej stosowanej do produkcji płyt PCB (<https://bit.ly/46Hmo7q>)

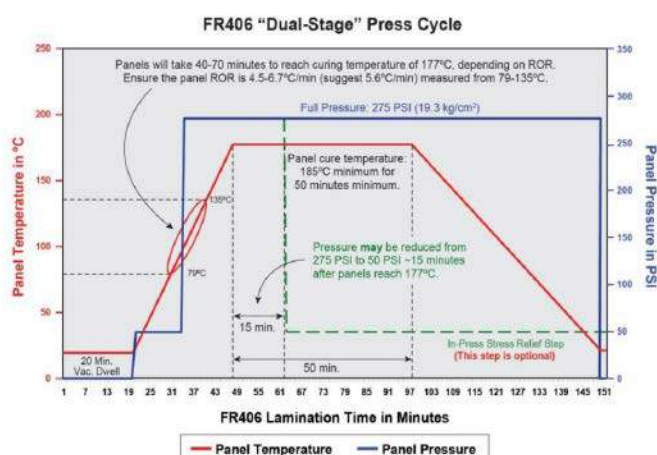
Oznakowanie folii	Najczęściej spotykane określenia	System metryczny		System calowy		
		Gramatura [g/m ²]	Grubość nominalna [μm]	Gramatura [oz./ft. ²]	Gramatura [g/254 in ²]	Grubość nominalna [mils]
E	5 μm	45,1	5,1	0,148	7,4	0,20
Q	9 μm	75,9	8,5	0,249	12,5	0,34
T	12 μm	106,8	12,0	0,350	17,5	0,47
H	1/2 oz, 18 μm	152,5	17,1	0,500	25,0	0,68
M	3/4 oz	228,8	25,7	0,750	37,5	1,01
1	1 oz, 35 μm	305,0	34,3	1	50,0	1,35
2	2 oz, 70 μm	610,0	68,6	2	100,0	2,70
3	3 oz	915,0	102,9	3	150,0	4,05
4	4 oz	1220,0	137,2	4	200,0	5,40
5	5 oz	1525,0	171,5	5	250,0	6,75
6	6 oz	1830,0	205,7	6	300,0	8,10
7	7 oz	2135,0	240,0	7	350,0	9,45
10	10 oz	3050,0	342,9	10	500,0	13,50
14	14 oz	4270,0	480,1	14	700,0	18,90

metody badania prepregów PCB w określonych warunkach temperatury i ciśnienia. W skrócie rzecz ujmując, metoda wg IPC TM-650 2.3.17.2 polega na pomiarze zmiany średnicy wewnętrznej okrągłych wycięt $\varnothing 1,000$ ", wykonanych w prepregu, poddawanemu następnie prasowaniu w ustalonych warunkach (**fotografia 8**). Dawniej jedyną

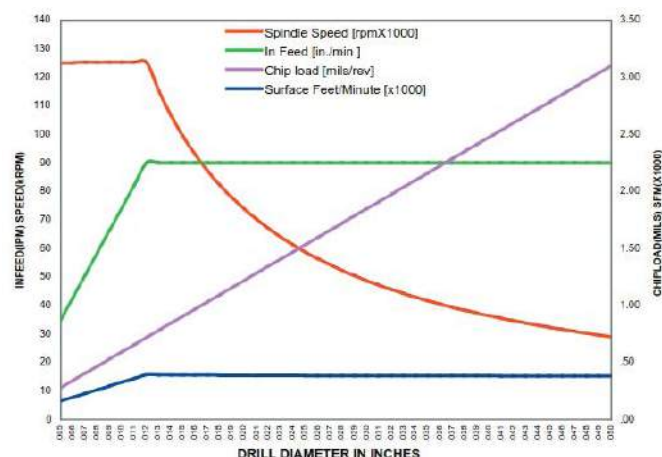
metodą pomiaru średnicy po prasowaniu było użycie suwmiarki – dziś można to wykonać numerycznie poprzez wyznaczenie obrysu otworu za pomocą oprogramowania, bazującego na wysokorozdzielczej fotografii arkusza testowego (**rysunek 6**). Nie zagłębiając się w szczegóły, powiemy tylko, że warto pamiętać o jednym, podstawowym fakcie – prepregi, określane dumnie przez producentów jako *low-flow* lub nawet *no-flow*, w rzeczywistości także „płyną”, choć w stopniu znacznie mniejszym niż materiały standardowe – porównanie lepkości dwóch wybranych prepregów pokazuje, że „niepłynące” prepregi w istocie wykazują lepkość o rząd wielkości wyższą niż wykonania klasyczne (**rysunek 7**).

Poszczególne prepregi różnią się ponadto pod względem wymogów procesu technologicznego, stosowanego do laminacji MLPCB. W każdym przypadku producent określa preferowane warunki prasowania, które na pierwszy rzut oka – pod względem przebiegu temperatury – mogą nieznacznie przypominać profile termiczne lutowania rozplwowego (**rysunki 8 i 9**). Producenci określają nawet szczegółowe warunki obróbki skrawaniem – przykładowy wykres parametrów wiercenia dla laminatów FR-4, produkowanych przez Taconic Synthane-Taylor, można zobaczyć na **rysunku 10**.

Warto wiedzieć, że nie wszystkie materiały są wzmacniane włóknem szklanym, różne pozostają także właściwości nowoczesnych prepregów w zakresie palności – podczas gdy producenci starają się dążyć do uzyskania klasy UL94V-0 (m.in. przez domieszkowanie bromem), to w niektórych przypadkach jest to niemożliwe, o czym rzecz jasna wytwórcy materiałów stosowanych do produkcji PCB muszą jasno poinformować potencjalnych odbiorców w notach katalogowych. Warto dodać, że ogólnościowy trend proekologiczny



Rysunek 9. Profil laminacji dwustopniowej dla laminatów FR406 marki ISOLA. Warto zwrócić uwagę na przebieg ciśnienia, które po początkowej wartości 50 psi (utrzymującej się przez kilkanaście minut) rośnie do 275 psi (<https://t.ly/m6EaQ>)



Rysunek 10. Parametry wiercenia dla laminatu FB650K marki Taconic Synthane-Taylor (https://t.ly/QSxM_)



Fotografia 9. Przykładowy arkusz typu bondply, dostępny w ofercie marki Rogers (<https://t.ly/PDVcc>)

Tabela 5. Szerokości ścieżek, dla których temperatura wzrośnie o 10°C przy przepływie prądu o podanej wartości; dane na podstawie IPC-2221 (<https://t.ly/hWA-F>)

Natężenie prądu [A]	Szerokość ścieżki [mm]			
	Miedź 35 µm		Miedź 70 µm	
	Warstwa zewnętrzna	Warstwa wewnętrzna	Warstwa zewnętrzna	Warstwa wewnętrzna
1	0,300	0,781	0,150	0,391
2	0,781	2,03	0,391	1,02
3	1,37	3,56	0,684	1,78
4	2,03	5,29	1,02	2,64
5	2,77	7,19	1,38	3,60
6	3,56	9,25	1,78	4,63
7	4,40	11,4	2,20	5,72
8	5,29	13,8	2,64	6,88
9	6,22	16,2	3,11	8,09
10	7,19	18,7	3,60	9,36
15	12,6	32,7	6,29	16,4
20	18,7	48,7	9,36	24,3
25	25,5	66,2	12,7	33,1
30	32,7	85,2	16,4	42,6

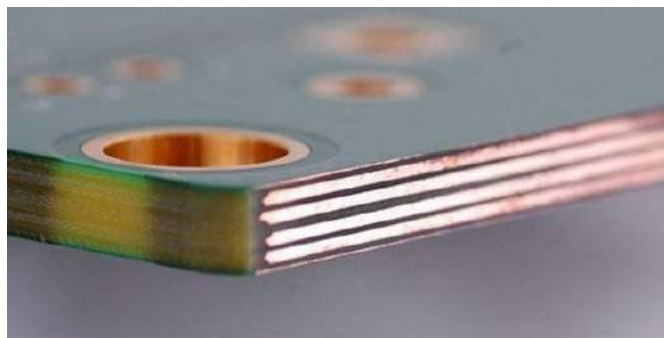
zawitał także do branży laminatów i prepregów, o czym świadczą bezhalogenkowe opóźniacze palenia, stosowane m.in. w laminatach z serii Terragreen marki ISOLA i spełniające wymogi normy IPC/JEDEC J-STD-709.

Kolejnym kryterium wyboru podczas opracowywania stosu warstw jest zakres materiałów, z którymi mogą być łączone poszczególne warstwy obwodu drukowanego. Niektóre materiały – np. poliimidowy prepreg drugiej generacji o oznaczeniu 38N marki Arlon – umożliwiają efektywne łączenie „trudnych” materiałów, takich jak podłoża kaptonowe, stosowane w obwodach sztywno-giętkich (*rigid-flex*) oraz w pełni elastycznych (*flex*). Z kolei słynna seria laminatów mikrofalowych RO4000 marki Rogers, wykonanych na bazie ceramiki węglowodorowej, wzmocnionej włókniną szklaną, zapewnia pełną kompatybilność z procesem technologicznym typowym dla FR-4, co znacząco zwiększa dostępność, także w przypadku producentów dysponujących parkiem maszynowym o ograniczonych możliwościach.

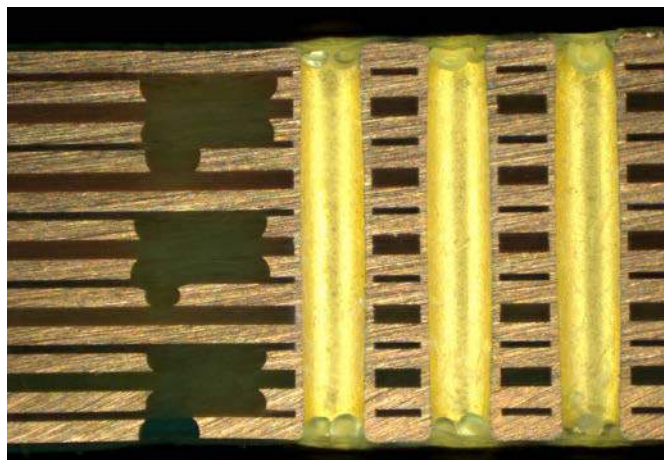
W szczególnie wymagających aplikacjach można także rozważyć stosowanie specjalnych arkuszy łączeniowych, określanych mianem *bondply* (fotografia 9) – firma Rogers oferuje szeroką gamę tego typu materiałów o grubości od 1,5 milsa, które pełnią funkcję „taśmy samoprzylepnej”, łączącej warstwy MLPCB wykonane z trudnych do połączenia materiałów, np. laminatów na bazie teflonu czy też... szkła.

Folia miedziana – cichy bohater w świecie obwodów drukowanych

Sporo miejsca poświęciliśmy do tej pory materiałom rdzeni i prepregów – a przecież w rzeczywistości głównym budulcem PCB, który ma (prawie) bezpośredni kontakt z komponentami, jest... miedź. Podstawowe oznaczenia oraz grubości i gramatury folii Cu, stosowanych do produkcji płyt wielowarstwowych (a także laminowanych fabrycznie na powierzchniach rdzeni), zebrano w tabeli 4. Zdecydowanie najpopularniejsze są rzecz jasna folie o grubości 18 µm, 35 µm oraz 70 µm, ale zdecydowanie nie należy uważać tej listy za wyczerpującą. Na rynku można bowiem znaleźć producentów, oferujących wytwarzanie PCB w technologii określonej jako *heavy copper*, czyli z użyciem warstw miedzi o ponadprzeciętnej grubości (fotografia 10). Tego typu rozwiązania – jak nietrudno się domyślić – mają zastosowanie przede wszystkim w obwodach dużej mocy, sprawdzają się ponadto w budowie transformatorów planarnych (fotografia 11). Oczywiście, nie ma nic za darmo – dostosowanie obróbki



Fotografia 10. Płyta 6-warstwowa, wykonana w technologii heavy copper (gramatura miedzi, zastosowanej do budowy warstw wewnętrznych, jest równa 8 oz) – <https://t.ly/V1SL5>



Fotografia 11. Przekrój 12-warstwowego transformatora planarnego dużej mocy, wykonanego na bazie miedzi o gramaturze 6 oz – całkowita grubość PCB wynosi 4 mm (<https://t.ly/B7cVt>)

tego typu płyt (nie wspominając o późniejszym ich lutowaniu) jest kluczowe dla poprawności procesu, przykładowo – parametry wiercenia są bardziej zbliżone do tych, które stosuje się w pracy z płytami miedzianymi, niż do standardowych warunków stosowanych przy obróbce PCB. Jako ciekawostkę warto natomiast dodać, że niektórzy producenci oferują możliwość wyprodukowania płyt, w których miedź ma... różną grubość w poszczególnych miejscach tej samej warstwy (fotografia 12), co niebawem zwiększa elastyczność projektowania i pozwala zmniejszyć wymiary płyt, w których na tej samej stronie montowane są zarówno elementy małej mocy (np. układy cyfrowe), jak i obwody przenoszące bardzo duże obciążenia.

W tabeli 5 pokazano szerokości ścieżek, dla których wzrost temperatury przy przepływie prądu o danej wartości wyniesie 10°C. Jak widać, izolacja termiczna, jaką stanowi laminat, znacząco redukuje amperaż ścieżek lub też – patrząc na to z drugiej strony – wymusza stosowanie szerszych połączeń dla obwodów zasilających czy

Tabela 6. Zalecane wartości minimalne odstępów i szerokości ścieżek dla różnych grubości miedzi na warstwach zewnętrznych (<https://t.ly/jnqOu>)

Gramatura miedzi na warstwie zewnętrznej [oz]	Minimalny odstęp izolacyjny [mil/mm]	Minimalna szerokość ścieżek [mil/mm]
0,33	3,50/0,089	3,00/0,08
0,50	4,00/0,10	3,50/0,089
1,00	6,00/0,15	5,50/0,14
2,00	9,00/0,23	9,00/0,23
3,00	13,00/0,33	13,00/0,33
4,00	17,00/0,43	17,00/0,43

Tabela 7. Wzrosty temperatury ścieżek o danej szerokości, spowodowane przepływem prądu o natężeniu od 0,4 A do 20,9 A (na podstawie wzorów z IPC-2221) – <https://t.ly/Hc6e1>

Wzrost temperatury →		10°	20°	30°	40°	50°	60°
Szerokość ścieżki							
[mm]	[mil]	Prąd [A]					
0,1	4	0,4	0,6	0,8	0,9	1,0	1,1
0,2	8	0,7	1,0	1,2	1,4	1,6	1,7
0,3	12	0,9	1,3	1,6	1,8	2,0	2,2
0,4	16	1,1	1,5	1,9	2,2	2,4	2,7
0,5	20	1,3	1,8	2,2	2,5	2,8	3,1
0,6	24	1,4	2,0	2,4	2,8	3,1	3,4
0,7	28	1,6	2,2	2,7	3,1	3,5	3,8
0,8	32	1,7	2,4	2,9	3,4	3,8	4,1
0,9	36	1,8	2,6	3,2	3,7	4,1	4,5
1,0	40	2,0	2,8	3,4	3,9	4,4	4,8
1,5	60	2,5	3,6	4,4	5,1	5,7	6,2
2,0	80	3,0	4,3	5,3	6,1	6,8	7,5
3,0	120	3,9	5,6	6,8	7,9	8,8	9,7
4,0	160	4,7	6,7	8,2	9,5	10,6	11,6
5,0	200	5,5	7,7	9,5	10,9	12,2	13,4
6,0	240	6,1	8,7	10,6	12,3	13,7	15,1
7,0	280	6,8	9,6	11,7	13,6	15,2	16,6
8,0	320	7,4	10,4	12,8	14,8	16,5	18,1
9,0	360	8,0	11,3	13,8	15,9	17,8	19,5
10,0	400	8,5	12,1	14,8	17,0	19,1	20,9

też sterujących elementami wykonawczymi. Dane pokazane w tabeli zostały opracowane na podstawie normy IPC-2221. Należy przy tym pamiętać, że grubość miedzi wpływa nie tylko na parametry prądowe i termiczne, ale także na dopuszczalne szerokości ścieżek i odstępy izolacyjne – im grubsza folia Cu, tym bardziej restrykcyjne ograniczenia dolnego zakresu tych wymiarów. Dla lepszej orientacji w tabeli 6 pokazano zalecane, minimalne szerokości ścieżek oraz odstępy izolacyjne dla sześciu najpopularniejszych grubości miedzi, zaś w tabeli 7 zebrano prognozowane wzrosty temperatury dla ścieżek o różnych szerokościach, przy zastosowaniu 35-mikrometrowej folii miedzianej. Warto dodać, że dane te dotyczą płyt dwuwarstwowych, zaś w przypadku konstrukcji MLPCB wewnętrzne płaszczyzny miedzi będą pełniły funkcję dodatkowego radiatora, przez co rzeczywiste wzrosty temperatury powinny być w praktyce niższe.

Wygodne w użyciu narzędzia graficzne, pozwalające na szybkie dokonanie estymacji bezpiecznej (choć często zawyżonej) granicy minimalnej szerokości ścieżek, znajdujemy w normie IPC-2221

(rysunek 11). Podążając w kierunku oznaczonym strzałką czerwoną, można z łatwością wyznaczyć amperaż (równy około 2,7 A) ścieżki o szerokości 140 milsów i grubości 35 μm , położonej na warstwie wewnętrznej (przy założeniu, że wzrost temperatury nie może przekroczyć 10°C). Z kolei strzałka pomarańczowa obrazuje sposób wyznaczania szerokości ścieżki (około 40 milsów) o grubości 18 μm , która – przewodząc prąd o natężeniu 1 A – może podgrzać się o nie więcej niż 30°C względem otaczających struktur płytki. W rzeczywistości dane te są przeszacowane, ale stanowią doskonały punkt wyjścia dla projektantów, którzy dbają o margines bezpieczeństwa podczas projektowania obwodów, zwłaszcza średniej i dużej mocy.

W przypadku obwodów RF problem doboru folii miedzianej staje się znacznie bardziej złożony – w praktyce okazuje się bowiem, że za straty w transmisji sygnałów odpowiada nawet... poziom chropowatości miedzi. Wynika to z prostego faktu – wraz ze wzrostem częstotliwości nasila się efekt naskórkowości, a zatem ścieżka o bardzo nierównej powierzchni będzie stanowiła większą przeszkodę dla



Fotografia 12. PCB w technologii heavy copper ze zróżnicowaną grubością miedzi w różnych obszarach tej samej warstwy (<https://t.ly/8lxxrk>)

REKLAMA

PROTOTYPOWA PRODUKCJA OBWODÓW DRUKOWANYCH PCB



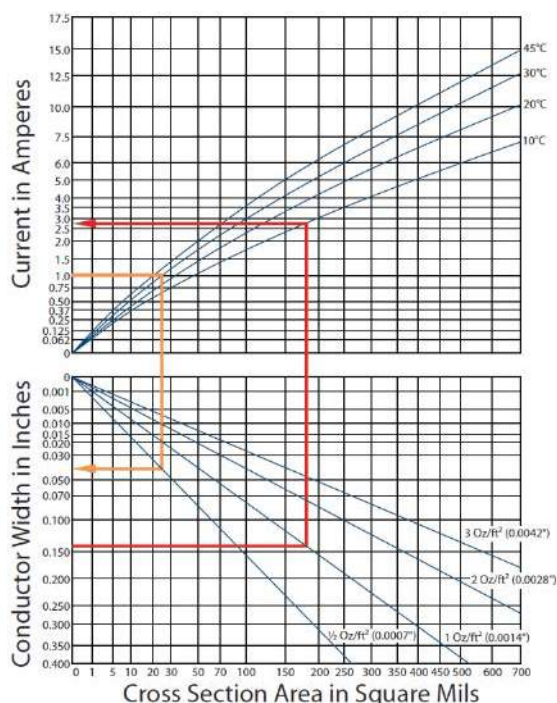
ul. Zeusa 61
80-299 Gdańsk
58 554 07 64
biuro@prototypy.com



- ekspresowa produkcja w 5h
- obwody wielowarstwowe
- prototypy PCB oraz serie
- laminat FR4, Rogers, MCPCB
- projekty układów elektronicznych
- montaż oraz uruchamianie prototypów
- precyzyjna produkcja z gwarancją jakości
- obsługa osób fizycznych i firm
- pojedyncza pełna płytka już od 25 zł netto
- inżynieria wsteczna
- wycena online

WWW.PROTOTYPY.COM





Rysunek 11. Nomogram wg IPC-2221, pozwalający na wyznaczenie szerokości ścieżki na warstwie wewnętrznej lub jej amperaży, przy znajomości pozostałych parametrów (grubości miedzi i dopuszczalnego wzrostu temperatury) (<https://t.ly/6F0Z7>)

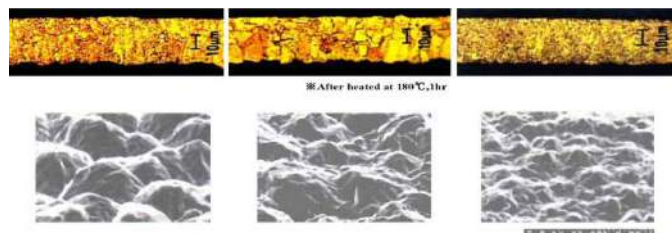
przepływu prądu niż linia transmisyjna o tych samych wymiarach i umieszczona na tym samym podłożu, ale o znacznie gładziej powierzchni. Przykładowe porównanie standardowej powierzchni folii Cu (wytwarzanej metodą osadzania elektrolitycznego) oraz folii typu *low-profile* (o obniżonej chropowatości) pokazano na **fotografii 13**. Mało tego – istnieje szereg klas chropowatości, obejmujących jeszcze szerszy zakres profili powierzchni (**fotografia 14**).

Tak duże zróżnicowanie sposobów wykończenia powierzchni folii miedzianej wynika z faktu, iż odpowiednie manipulowanie jej chropowatością pozwala dopasować rodzaj materiału do konkretnych potrzeb. Silne wiązanie z pozostałymi warstwami stosu PCB można uzyskać dzięki większej chropowatości – z tego też względu niektóre folie są dodatkowo matowione po stronie gładkiej (czyli tej, którą folia w trakcie produkcji była osadzona na bębnie). Z drugiej strony, obniżenie chropowatości pozwala wyraźnie poprawić transmisję sygnałów i obniżyć straty odbiciowe, co można zaobserwować na wynikach symulacji, pokazanych na **rysunku 12**.

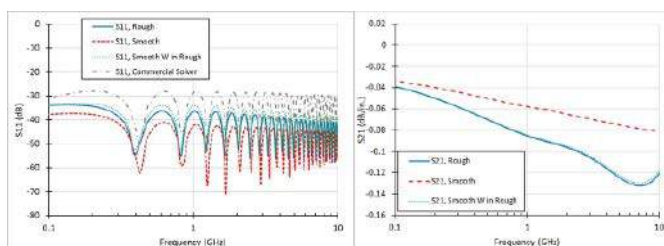


Matte Side of Standard Grade 1 Foil Matte Side of Low Profile Grade 1 Foil

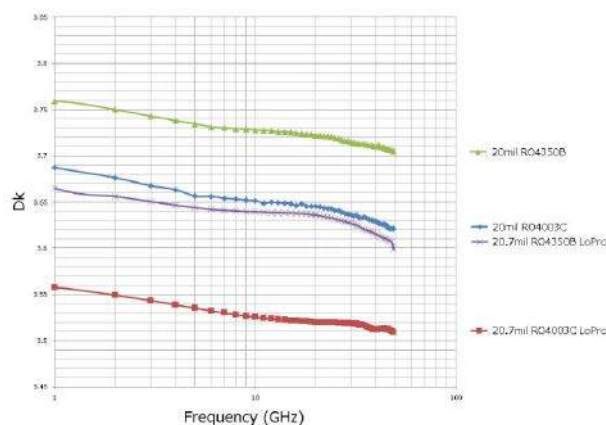
Fotografia 13. Porównanie chropowatości powierzchni folii miedzianych, wykonanych w wersji standardowej (po lewej) oraz niskoprofilowej (po prawej) – <https://t.ly/n0-vh>



Fotografia 14. Porównanie przekrojów i powierzchni folii miedzianej klasy (od lewej): HTE, Super-HTE oraz VLP (<https://t.ly/n0-vh>)



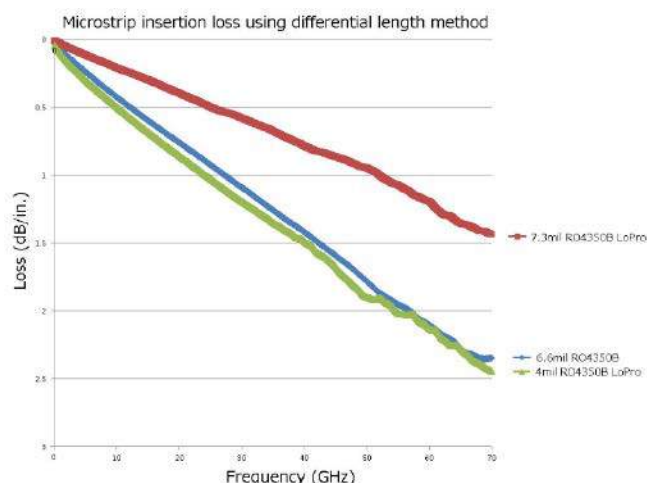
Rysunek 12. Symulacja parametrów S dla folii miedzianych o różnych poziomach chropowatości (<https://t.ly/UFBYm>)



Rysunek 13. Zależność stałej dielektrycznej od częstotliwości dla laminatów z serii RO4000 marki Rogers (https://t.ly/u_4ck)

Parametry elektryczne laminatów – czyli dlaczego na pytanie o wartość ϵ , FR-4 należy odpowiadać „to zależy”

W wielu opracowaniach, dotyczących rodzajów podłoży stosowanych do produkcji PCB, temat parametrów elektrycznych laminatów i prepregów bywa traktowany w pewnym sensie po macoszemu. Często powtarzana jest np. zdawkowa informacja, że w układach RF – zamiast klasycznych laminatów szklano-epoksydowych – warto stosować np. wysokiej klasy laminaty Rogers, które oferują lepsze parametry w zakresie stratności i stałej dielektrycznej. Tymczasem zarówno zgrubne przybliżenie stałych Dk dla FR-4, równe np. 4,5, jak i niepoparte konkretnymi wartościami wskazanie wyższości materiałów Rogers w zakresie RF to w istocie żadne informacje – dość powiedzieć, że w przypadku każdego materiału stała dielektryczna oraz stratność (parametr Df) nie tylko zmieniają się w funkcji częstotliwości przenoszonych sygnałów, ale także reagują na warunki środowiska (w tym temperaturę, a nawet wilgotność).



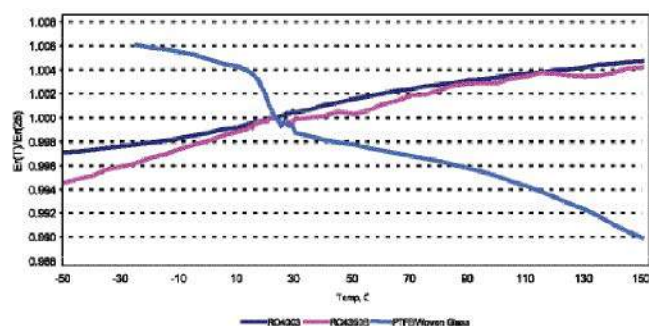
Rysunek 14. Zależność strat wtrąciennych (mierzonych w linii mikropaskowej) od częstotliwości dla laminatów z serii RO4000 marki Rogers (https://t.ly/u_4ck)

Tabela 8. Porównanie parametrów laminatów z serii FR408HR (<https://t.ly/jce15>)

Budowa rdzenia	Zawartość żywicy [%]	Grubość [cal]	Grubość [mm]	Stała dielektryczna (Dk)/Stratność (Df)					
				100 MHz	500 MHz	1 GHz	2 GHz	5 GHz	10 GHz
1×106	72,0%	0,0020	0,051	3,37/0,0075	3,36/0,0089	3,34/0,0096	3,32/0,0101	3,30/0,0107	3,30/0,0107
1×1080	57,0%	0,0025	0,064	3,67/0,0071	3,64/0,0079	3,62/0,0089	3,61/0,0092	3,60/0,0097	3,59/0,0095
1×1067	69,0%	0,0025	0,064	3,42/0,0075	3,40/0,0084	3,38/0,0095	3,36/0,0100	3,34/0,0105	3,33/0,0104
1×1080	63,0%	0,0030	0,076	3,54/0,0074	3,52/0,0082	3,50/0,0092	3,48/0,0096	3,47/0,0102	3,47/0,0101
1×1086	59,0%	0,0030	0,076	3,65/0,0072	3,63/0,0079	3,60/0,0091	3,59/0,0092	3,57/0,0098	3,57/0,0095
2×106	67,0%	0,0035	0,089	3,46/0,0074	3,45/0,0083	3,42/0,0094	3,40/0,0098	3,38/0,0104	3,37/0,0102
1×3313	51,0%	0,0035	0,089	3,82/0,0068	3,79/0,0076	3,77/0,0084	3,77/0,0087	3,74/0,0092	3,74/0,0090
1×106/1×1080	59,0%	0,0040	0,102	3,63/0,0072	3,61/0,0080	3,58/0,0090	3,57/0,0093	3,55/0,0098	3,54/0,0096
1×3313	55,0%	0,0040	0,102	3,72/0,0071	3,70/0,0077	3,68/0,0087	3,66/0,0090	3,65/0,0095	3,65/0,0094
1×106/1×1080	61,0%	0,0043	0,109	3,57/0,0073	3,56/0,0081	3,54/0,0092	3,52/0,0095	3,51/0,0099	3,50/0,0098
1×2116	50,0%	0,0045	0,114	3,82/0,0068	3,81/0,0075	3,80/0,0083	3,79/0,0086	3,78/0,0091	3,76/0,0089
1×106/1×1080	62,0%	0,0045	0,114	3,55/0,0073	3,54/0,0082	3,52/0,0092	3,50/0,0095	3,48/0,0100	3,48/0,0098
2×1067	69,0%	0,0050	0,127	3,42/0,0075	3,40/0,0084	3,38/0,0095	3,36/0,0100	3,34/0,0105	3,33/0,0104
2×1080	57,0%	0,0050	0,127	3,67/0,0071	3,64/0,0079	3,62/0,0089	3,61/0,0092	3,60/0,0097	3,59/0,0095
1×2116	54,0%	0,0050	0,127	3,72/0,0070	3,71/0,0077	3,70/0,0086	3,69/0,0089	3,67/0,0095	3,67/0,0093
2×1080	58,0%	0,0055	0,140	3,65/0,0072	3,63/0,0079	3,60/0,0091	3,59/0,0092	3,57/0,0098	3,57/0,0095
2×1086	58,0%	0,0060	0,152	3,65/0,0072	3,63/0,0079	3,60/0,0091	3,59/0,0092	3,57/0,0098	3,57/0,0095
1×1652	50,0%	0,0060	0,152	3,82/0,0068	3,81/0,0075	3,80/0,0083	3,79/0,0086	3,78/0,0091	3,76/0,0089
2×3313	50,0%	0,0070	0,178	3,82/0,0068	3,81/0,0075	3,80/0,0083	3,79/0,0086	3,78/0,0091	3,76/0,0089
2×3313	55,0%	0,0080	0,203	3,72/0,0071	3,70/0,0077	3,68/0,0087	3,66/0,0090	3,65/0,0095	3,65/0,0094
2×2116	54,0%	0,0100	0,254	3,72/0,0070	3,71/0,0077	3,70/0,0086	3,69/0,0089	3,67/0,0095	3,67/0,0093
2×1652	42,0%	0,0100	0,254	4,02/0,0065	4,01/0,0072	3,99/0,0080	3,98/0,0083	3,97/0,0088	3,96/0,0085
3×3313	55,0%	0,0120	0,305	3,72/0,0071	3,70/0,0077	3,68/0,0087	3,66/0,0090	3,65/0,0095	3,65/0,0094
1×1080/2×1652	49,0%	0,0140	0,356	3,86/0,0068	3,83/0,0074	3,81/0,0083	3,80/0,0086	3,78/0,0090	3,78/0,0089
3×1652	45,0%	0,0160	0,406	3,95/0,0066	3,92/0,0072	3,91/0,0080	3,90/0,0083	3,89/0,0086	3,88/0,0085
3×1652	50,0%	0,0180	0,457	3,82/0,0068	3,81/0,0075	3,80/0,0083	3,79/0,0086	3,78/0,0091	3,76/0,0089
2×2116/2×1652	46,0%	0,0190	0,483	3,92/0,0066	3,90/0,0073	3,89/0,0081	3,87/0,0089	3,86/0,0087	3,85/0,0086
2×2116/2×1652	47,0%	0,0200	0,508	3,89/0,0067	3,88/0,0074	3,87/0,0083	3,85/0,0085	3,84/0,0089	3,83/0,0088
2×2116/2×1652	50,0%	0,0210	0,533	3,82/0,0068	3,81/0,0075	3,80/0,0083	3,79/0,0086	3,78/0,0091	3,76/0,0089
2×2116/3×1652	52,0%	0,0280	0,711	3,77/0,0073	3,76/0,0080	3,76/0,0088	3,75/0,0091	3,73/0,0093	3,73/0,0093

Mało tego – wyjściowe wartości zależą w dużej mierze nawet od zastosowanej konstrukcji laminatu. Producenci rdzeni FR-4 bazują na prepegach, które – składane razem w celu uzyskania odpowiedniej grubości wynikowej – wpływają na osiągane parametry wysokoczęstotliwościowe. W tabeli 8 znajdują się przykładowe dane dla laminatów z serii FR408HR marki ISOLA, zaś tabela 9 prezentuje analogiczne wartości dla samych prepegów – a wciąż mówimy tutaj tylko o jednej, konkretnej rodzinie materiałów z oferty jednego producenta!

W przypadku porównywania podłoży FR-4 z materiałami wykonanymi na bazie ceramiki, teflonu czy też poliimidu zauważymy, że rozrzut wartości Dk, Df (jak i wszelkich pozostałych parametrów) jest już zdecydowanie większy – w tabeli 10 zaprezentowano zestawienie 14 różnych materiałów wraz z wybranymi wartościami Dk, Df (przy danej częstotliwości testowej), a także kilka parametrów termicznych, którym więcej uwagi poświęcimy w dalszej części niniejszego artykułu.



Rysunek 15. Wykres względnych zmian stałej dielektrycznej w zależności od częstotliwości dla laminatów z serii RO4000 marki Rogers (linia granatowa – RO4003, różowa – RO4350B) w porównaniu ze stosowaną charakterystyką dla PTFE wzmacnianego włókniną szklaną. Wartości na osi pionowej odnoszą się do stałej zmierzonej w temperaturze 25°C (https://t.ly/u_4ck)

REKLAMA



**obwody drukowane
jedno, dwustronne
oraz wielowarstwowe
- obwody drukowane
do zastosowań LED
- prototypy i serie
produkcyjne**

tel. +48 12 636 00 33
tel. +48 604 773 095
www.hatron.com
hatron@hatron.com

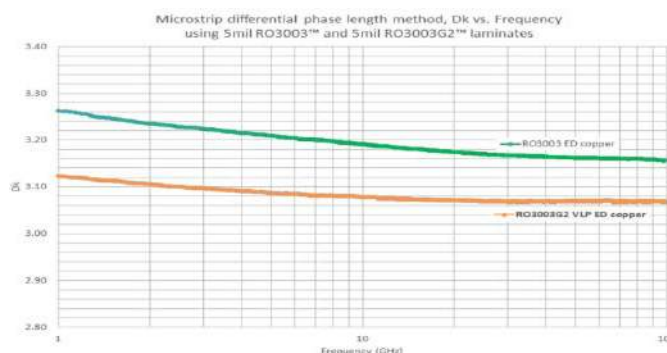


Tabela 9. Porównanie parametrów prepregów z serii FR408HR (<https://t.ly/jce15>)

Prepreg	Zawartość żywicy [%]	Grubość [cal]	Grubość [mm]	Stała dielektryczna (Dk)/Stratność (Df)					
				100 MHz	500 MHz	1 GHz	2 GHz	5 GHz	10 GHz
106	70,0%	0,0019	0,048	3,40/0,0080	3,38/0,0090	3,36/0,0100	3,34/0,0100	3,32/0,0110	3,31/0,0110
106	75,0%	0,0023	0,058	3,30/0,0080	3,27/0,0090	3,26/0,0100	3,23/0,0110	3,22/0,0110	3,22/0,0110
1067	69,0%	0,0023	0,058	3,42/0,0080	3,40/0,0080	3,38/0,0100	3,36/0,0100	3,34/0,0110	3,33/0,0100
1067	73,0%	0,0027	0,069	3,33/0,0080	3,32/0,0090	3,30/0,0100	3,28/0,0100	3,26/0,0110	3,26/0,0110
1086	62,0%	0,0032	0,081	3,55/0,0070	3,54/0,0080	3,52/0,0090	3,50/0,0100	3,48/0,0100	3,48/0,0100
1080	65,0%	0,0032	0,081	3,50/0,0070	3,49/0,0080	3,46/0,0090	3,44/0,0100	3,42/0,0100	3,42/0,0100
1086	65,0%	0,0035	0,089	3,50/0,0070	3,49/0,0080	3,46/0,0090	3,44/0,0100	3,42/0,0100	3,42/0,0100
1080	71,0%	0,0040	0,102	3,37/0,0080	3,36/0,0090	3,34/0,0100	3,32/0,0100	3,30/0,0110	3,30/0,0110
2313	56,0%	0,0040	0,102	3,70/0,0070	3,67/0,0080	3,65/0,0090	3,64/0,0090	3,63/0,0100	3,62/0,0100
3313	57,0%	0,0040	0,102	3,67/0,0070	3,64/0,0080	3,62/0,0090	3,61/0,0090	3,60/0,0100	3,59/0,0100
2113	57,5%	0,0041	0,104	3,67/0,0070	3,64/0,0080	3,62/0,0090	3,61/0,0090	3,60/0,0100	3,59/0,0100
3313	59,0%	0,0042	0,107	3,63/0,0070	3,61/0,0080	3,58/0,0090	3,57/0,0090	3,55/0,0100	3,54/0,0100
2116	55,0%	0,0049	0,124	3,72/0,0070	3,70/0,0080	3,68/0,0090	3,66/0,0090	3,65/0,0100	3,65/0,0090
2116	58,0%	0,0054	0,137	3,65/0,0070	3,63/0,0080	3,60/0,0090	3,59/0,0090	3,57/0,0100	3,56/0,0100

Aby jeszcze lepiej zrozumieć właściwości nowoczesnych laminatów mikrofalowych (których przykłady także trafiły do tabeli 10), warto przestudiować wykresy pokazane na **rysunkach 13...15** – jak widać, zachowanie tych materiałów jest dalece stabilniejsze w funkcji częstotliwości niż dla markowych laminatów FR-4 – przykładowo, procentowa zmiana stałej Dk dla RO4000, w bardzo szerokim paśmie 1...50 GHz, jest porównywalna z analogiczną zmianą tego parametru dla FR408HR, ale w 5-krotnie węższym przedziale częstotliwości! Zachowanie opisywanych laminatów Rogers jest ponadto znacznie bardziej „przewidywalne”, niż w przypadku innych materiałów, np. bazujących na PTFE wzmacnianym włókniną szklaną – stosowne wykresy można zobaczyć na rysunku 15.

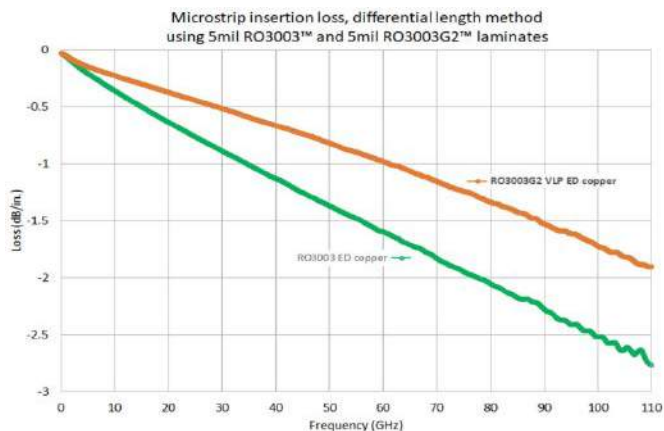
Inne rodzaje laminatów oferują jeszcze lepsze parametry w zakresie zarówno stabilności stałej dielektrycznej, jak i strat przy przenoszeniu sygnałów w paśmie mikrofalowym. Dane dla serii Rogers RO3003G2 można zobaczyć na **rysunkach 16 i 17** – warto zwrócić uwagę, jak doskonale widać tutaj wpływ rodzaju zastosowanej miedzi i to zarówno na Df, jak i Dk. Wykresy oznaczone zieloną linią odnoszą się do laminatów ze standardową folią miedzianą, wykonaną

Rysunek 16. Zależność stałej dielektrycznej od częstotliwości dla laminatów z serii RO3003 marki Rogers (<https://t.ly/otuMZ>)

w technologii ED (tj. przez elektroosadzanie), zaś wykresy pomarańczowe to charakterystyki tego samego laminatu, ale obłożonego folią o bardzo niskim profilu chropowatości powierzchni (VLP ED od *very low profile ED*). Jak widać, różnica w stratności przekracza 0,5 dB/

Tabela 10. Wybrane właściwości niektórych rodzajów laminatu i prepregów (<https://bit.ly/3s61IXE>)

Materiał	Stała dielektryczna Dk @ f _{TEST}	Stratność Df @ f _{TEST}	Częstotliwość testu f _{TEST}	Wsp. rozszerzalności termicznej CTE X;Y;Z	Temperatura zeszklenia Tg [°C]	Temperatura dekompozycji Td [°C]	Absorpcja wody [%]
FR-4 TG 135	4,6	0,015	1 MHz	Z 55	140	310	0,15
FR-4 TG 150	4,5	0,016	1 MHz	Z 52	155	335	0,09
FR-4 TG 170	4,8	0,013	1 MHz	Z 45	180	345	0,10
BT-Epoxy	4,1	0,01	1 GHz	14;14;55	180	325	–
Rogers 4350	3,48	0,004	10 GHz	14;16;50	280	390	0,06
Rogers 4003	3,38	0,0027	10 GHz	11;14;46	280	425	0,06
Rogers 4403 Prepreg	3,17	0,005	10 GHz	16;19;80	280	390	0,10
Rogers 4450 Prepreg	3,54	0,004	10 GHz	19;17;60	280	390	0,10
Taconic RF 35	3,5	0,0018	10 GHz	21;21;64	–	–	0,02
Arlon DiClad 522 Teflon	2,65	0,0022	10 GHz	14;21;173	–	–	0,03
Arlon DiClad 527 Teflon	2,65	0,0022	10 GHz	14;21;173	–	–	0,03
Arlon 85NT Polyimide	3,7	0,015	1 GHz	9; 9;90	280	390	0,60
Bergquist HPL03015	6,6	0,005	1 MHz	–	185	–	3,0
Megatron 6 Panasonic	3,6	0,002	2 GHz	–	210	410	–



Rysunek 17. Zależność strat wtrąceniowych w linii mikropaskowej od częstotliwości – dane dla laminatów z serii RO3003 marki Rogers (<https://t.ly/otuMZ>)

cał już dla częstotliwości 50 GHz i niemal jednostajnie rośnie wraz ze wzrostem pasma.

Pozostałe parametry materiałowe PCB

Stała dielektryczna i stratność, a także ich zależność od częstotliwości czy temperatury to parametry zdecydowanie najczęściej używane do porównywania laminatów, głównie (choć nie tylko) w zakresie aplikacji RF oraz szybkich systemów cyfrowych. Nie należy jednak zapominać, że w praktyce duże znaczenie ma także szereg innych wielkości, charakteryzujących rdzenie oraz prepregi. Poniżej opisujemy pokrótce najważniejsze z nich, w nawiasach okrągłych podając przykładowe wartości, zaczerpnięte z noty katalogowej wspomnianej już wcześniej rodziny FR408HR (<https://t.ly/tIYo4>).

- **Temperatura zeszklenia, T_g (190°C)** określa punkt termiczny, w którym następuje zmiana struktury materiału ze sztywnego na plastyczny (czy też „gumopodobny”), a wiąże się z tym także dość gwałtowna zmiana pojemności cieplnej materiału oraz jego rozszerzalności termicznej. Materiały o wysokiej wartości T_g oferują niższą wartość współczynnika rozszerzalności cieplnej, lepszą odporność na stropy termiczne podczas lutowania oraz późniejszej pracy (w tym cyklicznej), a także mniejszą podatność na delaminację oraz... niższy poziom absorpcji wilgoci.
- **Współczynnik rozszerzalności cieplnej, CTE (XY: 16 ppm/°C, Z: 55 ppm/°C)**, podobnie jak w przypadku innych rodzajów materiału, określa stopień, w jakim próbka rozszerza się w miarę wzrostu temperatury. Parametr ten ma ogromne znaczenie dla niezawodności urządzeń narażonych na liczne cykle zmian temperatury, bowiem powtarzalne ogrzewanie i chłodzenie laminatu może prowadzić do powstawania mikrouszkodzeń. Wartość CTE dla materiałów anizotropowych jest podawana osobno dla osi X, Y i Z lub przynajmniej w postaci dwóch wartości (jedna dla płaszczyzny XY i druga dla osi Z, prostopadłej do płaszczyzny laminatu). Z uwagi na stabilność wymiarową włókna szklanego w funkcji temperatury, laminaty na bazie włókna szklanego rozszerzają się podczas ogrzewania zdecydowanie wolniej w płaszczyźnie laminatu (XY) niż w kierunku prostopadłym (Z) – co więcej, w miarę zbliżania się do punktu T_g , rozszerzalność w kierunku Z rośnie i to drastycznie – w przypadku FR408HR, z początkowych 55 ppm/°C, aż do 230 ppm/°C!
- **Temperatura dekompozycji, T_d (360°C)** – w tym punkcie niektóre związki chemiczne zawarte w żywicy zaczynają ulegać rozkładowi, co wiąże się z utratą masy oryginalnej próbki. Dokładne określenie T_d jest możliwe tylko wtedy, gdy podany zostanie także procentowy próg (5%) utraty masy materiału. Dekompozycja powoduje rzecz jasna zmianę właściwości fizykochemicznych żywicy, a co za tym idzie – nieodwracalne uszkodzenie laminatu.
- **Czas do delaminacji (T_{260} : 60 minut, T_{288} : > 30 minut)** to czas, po którym (w danej temperaturze) rozpoczyna się proces rozwarstwienia

struktury materiału; innymi słowy, w objętości laminatu zaczynają pojawiać się luki, zdolne do dokonania licznych uszkodzeń w strukturze płytki drukowanej (przede wszystkim w rejonie padów i przelotek).

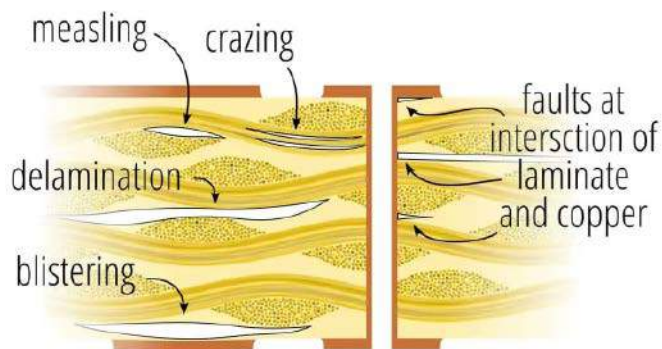
- **Przewodność cieplna (0,4 W/m·K)** – w przypadku większości laminatów (z wyjątkiem podłoży aluminiowych oraz materiałów specjalnego przeznaczenia) przewodność cieplna jest dość niska, co przekłada się m.in. na wspomniane wcześniej problemy z odprowadzaniem nadmiaru ciepła z warstw wewnętrznych. Istnieją jednak specjalistyczne odmiany podłoży, które oferują wielokrotnie wyższą przewodność, a co za tym idzie – lepiej nadają się do budowy m.in. urządzeń mikrofalowych, wymagających efektywnego chłodzenia komponentów o zwiększonej mocy strat. Przykładowo – seria TC350 marki Rogers charakteryzuje się przewodnością cieplną na poziomie 1,24 W/m·K, czyli ponad trzykrotnie wyższą w porównaniu do FR-4.
- **Rezystywność objętościowa i powierzchniowa** (odpowiednio – $4,4...9,4 \times 10^7 \text{ M}\Omega/\text{cm}$, $0,026...2,1 \times 10^8 \text{ M}\Omega$) to parametry definiowane w określonych warunkach środowiskowych, tj. po ekspozycji na wilgoć lub podwyższoną temperaturę. Duże znaczenie mają zwłaszcza w urządzeniach wysokonapięciowych oraz precyzyjnej aparaturze pomiarowej, dla której nawet minimalne upływności mogą stanowić istotne zaburzenie uzyskiwanych wyników.
- **Wytrzymałość elektryczna (1741 V/mil, czyli 70 kV/mm)** określa odporność rdzenia (bądź prepregu po laminacji) na wysokie napięcie, przyłożone do próbki po obu jej stronach. W notach katalogowych można ponadto natrafić na progową wartość napięcia przebicia, wyrażaną w kV (w tym przypadku: > 50 kV).
- **Odporność na łuk elektryczny (137 s)** – określa czas, po którym na powierzchni próbki może pojawić się łuk elektryczny, co ważne – w warunkach wysokiego napięcia, ale ograniczonego natężenia prądu.
- **Siła odrywania (0,9...1,14 N/mm)** – parametr opisujący siłę, niezbędną do zerwania paska folii miedzianej z powierzchni laminatu pod kątem prostym do powierzchni podłoża i w ściśle określonych warunkach laboratoryjnych, w tym po wywołaniu „stresu termicznego”. Jednostka opisywanej wielkości odwołuje się do szerokości testowanej ścieżki.
- **Absorpcja wilgoci (0,061%)** – parametr ten opisuje procentową zmianę masy próbki po poddaniu jej „namaczaniu” w określonych normą warunkach. Dla laminatów FR-4 ma wartość oscylującą przeważnie około 0,05%, ale dokładny wynik zależy w dużej mierze od opisanych wcześniej parametrów spłotu włókna szklanego, żywicy, zastosowanych wypełniaczy, etc.

REKLAMA

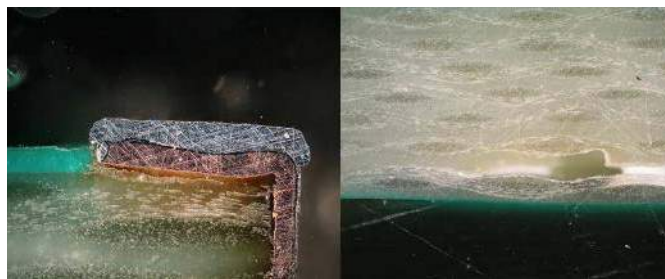
ELMAX
1988

OBWODY DRUKOWANE
Produkcja, Projektowanie, Montaż

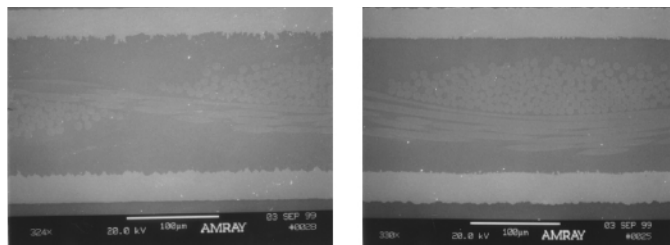
Certyfikat Underwriters Laboratories 94V-0 E480148 TYPE 1 Zakład produkcyjny: 05-660 Warka ul. M. Ropielewskiej 17 tel. 22 781 63 95 22 761 95 80 fax. 22 781 63 95 w 23 www.elmax.waw.pl elmax@elmax.waw.pl	Płytki jednostronne Płytki dwustronne Płytki na podłożu aluminium Płyty czołowe FR4	Serie dowolne Prototypy Maksymalny wymiar płytek 1w 630 mm
	Dokumentacja technologiczna Dokumentacja konstrukcyjna Trawione szablony SMD	Montaż elektroniczny Krótkie terminy Wykonania super expresowe
	Aktywny kalkulator prototypów na stronie internetowej	Pokrycie Sn lub SnPb inne na życzenie Maski, opisy montażowe w różnych kolorach



Rysunek 18. Rodzaje uszkodzeń materiałowych w strukturze laminatów PCB (https://t.ly/s_jQa)



Fotografia 16. Widoki przekrojów PCB, obrazujące przykładowe uszkodzenia: po lewej – uniesienie pierścienia padu metalizowanego (PTH) i oderwanie go od powierzchni laminatu; po prawej – wnęka wewnątrz materiału, spowodowana delaminacją (https://t.ly/s_jQa)



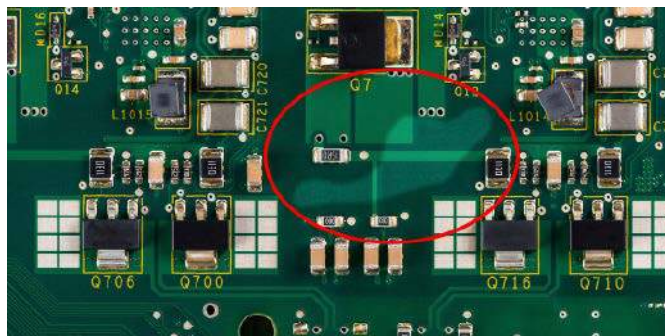
Fotografia 15. Przekroje laminatu z folią standardową (po lewej) oraz z dodatkowym matowaniem po stronie bębna (DST – drum side treatment) – po prawej. Źródło: <https://t.ly/0ytQY>

Wybrane defekty materiałowe w płytach wielowarstwowych, czyli koszmar senny elektronika

Na koniec pozostawiliśmy temat, z którym żaden praktykujący elektronik z pewnością nie chciałby zetknąć się w praktyce. Niestety, trudne warunki, jakim poddawane są płytki drukowane zarówno w procesach montażowych (zwłaszcza podczas lutowania rozpliwowego lub na fali), jak i w późniejszej eksploatacji, narażają PCB na liczne rodzaje uszkodzeń. Co gorsza, najbardziej podatne są właśnie płyty wielowarstwowe, w szczególności te o największej gęstości upakowania (HDI) – niewielkie przelotki i pady elementów THT idą zwykle na przysłowiowy pierwszy ogień, a wynika to przede wszystkim z pracy materiałów podczas cykli termicznych.

Podstawowe rodzaje uszkodzeń, wynikające głównie ze struktury laminatów, pokazano schematycznie na **rysunku 18**.

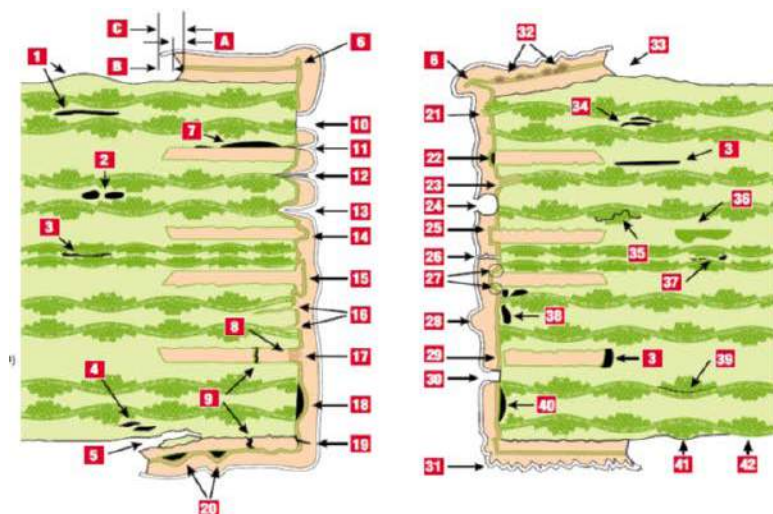
- **Measling** to mikrouszkodzenie, przebiegające na styku włókien szklanych wątku i osnowy. Z zewnątrz (nawet przez warstwę soldermaski) może być to widoczne jako niewielkie przejaśnienia



Fotografia 17. Delaminacja widoczna jako „przebarwienie soldermaski” (<https://t.ly/zDHQP>)

w materiale laminatu, zaś w rzeczywistości to nic innego, jak pusta wnęka, wybudowana pomiędzy przecinającymi się prostopadle włóknami. Niektóre standardy dopuszczają istnienie zjawiska *measlingu*, o ile nie mamy do czynienia z najwyższej klasy PCB do celów wojskowych lub lotniczych albo z urządzeniem wysokonapięciowym.

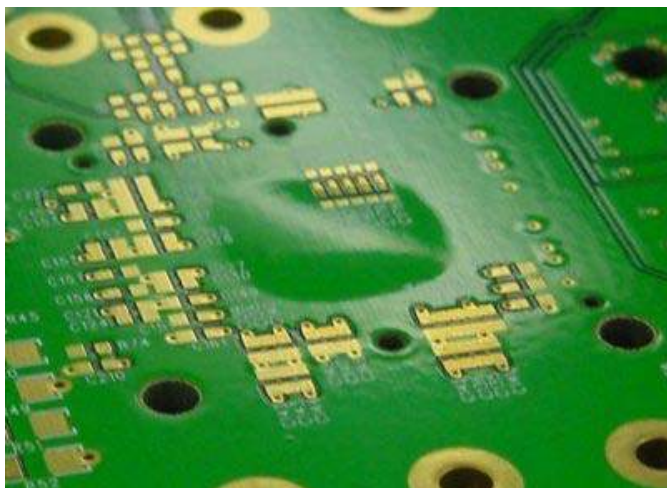
- **Crazing** – zjawisko polegające na odseparowaniu od siebie włókien szklanych, widoczne zwykle jako niewielkie, punktowe przejaśnienia lub skrzyżowane linie, odwzorowujące układ splotu włókniny. Do tego typu uszkodzeń dochodzi zwykle w wyniku przeciążeń mechanicznych.
- **Delaminacja** to rodzaj relatywnie obszernego uszkodzenia, będącego efektem wytworzenia płaskiej wnęki pomiędzy rdzeniem a folią miedzianą bądź dowolnymi innymi arkuszami stosu. Z uwagi na potencjalnie poważne konsekwencje dla niezawodności, delaminacja jest zwykle niedopuszczalna, choć wg IPC-A-610 może obejmować nie więcej niż 25% odległości



- | | |
|-----------------------------|-----------------------------|
| A Undercut | 21 Glass Fiber Prostitution |
| B Outgrowth | 22 D-Effect |
| C Overhang | 23 Wicking |
| 1 (Resin) Blistering | 24 Void (Metal Resist) |
| 2 Laminate Void | 25 (Positive) Etchback |
| 3 (Resin) Delamination | 26 Barrel Crack |
| 4 Lifted Land Crack | 27 Shadowing |
| 5 Pad Lifting | 28 Nodule |
| 6 Burr | 29 Rest Smear (ICD) |
| 7 Pink Ring | 30 Void (Resist Residues) |
| 8 Negative Etchback | 31 Burned Plating |
| 9 Foil Crack | 32 Starburst |
| 10 Void (PTH) | 33 Pad Rotation |
| 11 Wedge Void | 34 Resin Crack |
| 12 Glass Void | 35 Crazing |
| 13 Microvoid (Glass) | 36 Foreign Inclusion |
| 14 Arrow Heading | 37 Prepreg Void |
| 15 Nail Heading | 38 Pocket Void |
| 16 Drilling Cracks | 39 Measling |
| 17 Innerlayer Burning (ICD) | 40 Resin Recession |
| 18 Pull Away | 41 Glass-Weave Texture |
| 19 Corner Crack | 42 Glass-Weave Exposure |
| 20 Blistering | |

Originally designed by
Vladyslav Mommers BV, Netherlands
Reviewed by STPPTH
Rotech Deutschland GmbH, Berlin

Rysunek 19. Możliwe rodzaje uszkodzeń w strukturze padu PTH lub przelotki (<https://t.ly/pOcNA>)



Fotografia 18. Delaminacja przebiegająca z widocznym uniesieniem powierzchni warstwy zewnętrznej PCB (<https://t.ly/OFG8j>)

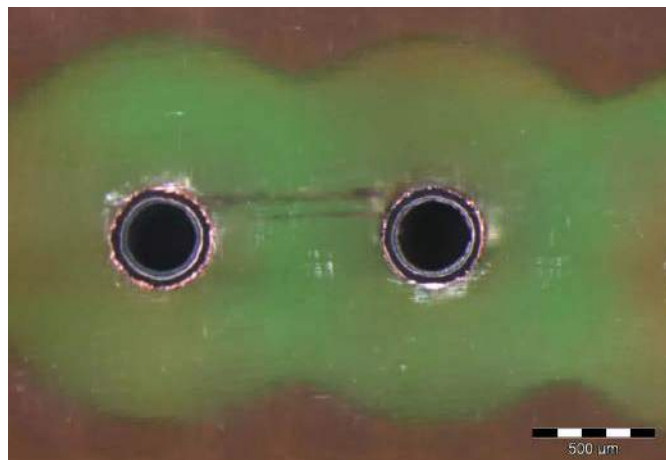
między padami THT lub ścieżkami na warstwach wewnętrznych. Najczęstszą przyczyną delaminacji są stropy termiczne.

- **Blistering** – odmiana delaminacji, przebiegająca najbardziej powierzchniowo, z widocznymi pęcherzami powietrza, powodującymi uniesienie warstwy zewnętrznej folii miedzianej wraz z soldermaską. Stanowi istotne zagrożenie dla niezawodności PCB, gdyż w oczywisty sposób grozi uszkodzeniem ścieżek, padów, a nawet komponentów.

Dwa przykłady uszkodzeń, uwidocznione na przekrojach PCB, można zobaczyć na **fotografii 16**.

Wspomnieliśmy już, że najbardziej podatne na uszkodzenia są przelotki (zwłaszcza te przebiegające przez cały stos PCB) oraz pady elementów przewlekanych. Co gorsza, w przypadku płyt wielowarstwowych wiele z tych uszkodzeń przebiega na styku warstw wewnętrznych oraz metalizacji otworu – mrozące krew w żyłach każdego elektronika podsumowanie znanych rodzajów defektów można zobaczyć na **rysunku 19**. Należy dodać, że zdecydowana większość z nich to w istocie pokłosie błędów w produkcji – np. nieprawidłowego przechowywania lub braku właściwego kondycjonowania materiałów składowych (głównie laminatów i prepregów) przed rozpoczęciem produkcji płyty wielowarstwowej. Rzecz jasna, wiele awarii może także pojawić się w późniejszym okresie „życia” PCB – mowa chociażby o delaminacji i *blisteringu* (dwa takie przypadki pokazano na **fotografiach 17 i 18**).

Istotne ryzyko niesie ze sobą jeszcze jeden, bardzo poważny rodzaj uszkodzeń, określany mianem CEF (od *conductive anodic filament*, czyli wzrost przewodzącego włókna). Efekt ten wynika ze zjawiska elektromigracji – zanieczyszczenia jonowe, obecne w strukturze laminatu, mogą w dłuższej perspektywie migrować pomiędzy blisko położonymi przewodnikami (np. otworami metalizowanymi lub ścieżkami), znajdującymi się pod napięciem stałym, prowadząc w efekcie do powstania mikroskopijnej zwory (**fotografia 19**). Warto dodać, że choć „preferowana” lokalizacja dla tworzenia się tego typu filamentów przebiega wzdłuż włókien szklanych (bowiem właśnie wtedy najłatwiej zachodzi migracja jonów), to wzrost włókien przewodzących może być także efektem wtórnym do powstałych uprzednio wnęk delaminacyjnych. Co więcej,



Fotografia 19. Zwarcie spowodowane wzrostem przewodzącego włókna miedzi pomiędzy padami, znajdującymi się pod napięciem stałym (<https://t.ly/-c-Dh>)

włókno nie musi wcale bezpośrednio połączyć dwóch przewodników – przyczyną awarii może być po prostu... znaczący spadek rezystancji, gdy dystalne zakończenie włókna znajdzie się w niewielkiej odległości od przeciwległej ścieżki bądź otworu.

Wykrycie dokładnego położenia włókien CAF w głębszych rejonach PCB okazuje się bardzo trudne lub wręcz niemożliwe, przynajmniej przy ograniczeniu się do badań nieniszczących. Dlatego też bardzo duże znaczenie ma zapewnienie odpowiednich warunków przechowywania i obróbki (laminowania) PCB, warto też wiedzieć, że prawdopodobieństwo wystąpienia CAF jest nieporównanie wyższe w przypadku laminatów bazujących na splocie standardowym niż dla spłotu rozszerzonego mechanicznie (MS).

Podsumowanie

W artykule pokazaliśmy szereg istotnych, a nader często pomijanych (lub przynajmniej traktowanych po macoszemu) zagadnień materiałowych, związanych z produkcją płyt wielowarstwowych. Nietrudno dojść do wniosku, że prawidłowe projektowanie obwodów MLPCB, zwłaszcza w przypadku aplikacji, w których oczekujemy najwyższej niezawodności oraz doskonałych parametrów sygnałowych, wymaga od projektanta poświęcenia uwagi zagadnieniom budowy stosu – i to nie tylko w zakresie doboru liczby i rozmieszczenia warstw, ale także grubości poszczególnych rdzeni i prepregów oraz skrupulatnego doboru zastosowanych materiałów. Współczesny przemysł dysponuje setkami odmian laminatów, prepregów oraz folii miedzianych, które są w stanie sprostać najbardziej wyśrubowanym wymaganiom – osobnym tematem pozostaje natomiast ich dostępność w magazynach wytwórców PCB. Mimo wszystko, w myśl dawnego powiedzenia: „koniec języka za przewodnika”, przed zleceniem produkcji warto skontaktować się z działem technologicznym preferowanego producenta obwodów drukowanych i dopytać o posiadane możliwości doboru materiałów. A najlepiej jest zrobić to na wczesnym etapie projektu, kiedy wieloma parametrami można jeszcze dość dowolnie żonglować w celu uzyskania optymalnych osiągnięć urządzenia.

inż. Przemysław Musz, EP

REKLAMA

EP.com.pl

Odwiedź stronę z mnóstwem doskonałych projektów

Wielowarstwowe PCB, od prostych do awangardowych

Każda płytką PCB to coś wyjątkowego, niepowtarzalnego – Twoje dzieło. Wykonamy je dla Ciebie z najwyższą starannością. W technologii BASIC realizujemy płytki PCB jednostronne, dwustronne i wielowarstwowe. Nasze standardy stanowią podstawę jakości i wydajności w postaci sprawdzonych stosów i zasad projektowania. Dzięki standardom zapobiegasz błędom i zwiększasz produktywność w projektowaniu, wykonalności i niezawodności PCB w jej aplikacji.

W jaki sposób Twój projekt staje się płytką drukowaną? Nasze standardowe stosy wielowarstwowych obwodów BASIC są produkowane poprzez laminowanie folii. W tym celu wytrawione rdzenie warstw wewnętrznych układa się w stosy z prepregami i folią miedzianą, a następnie prasuje. Laminacja folią,

w porównaniu do laminacji rdzeni, ma następujące zalety:

- tylko jeden rdzeń musi zostać wytrawiony,
- nie ma potrzeby zabezpieczania jednej strony podczas procesu trawienia,
- nie jest konieczne dopasowywanie folii do rdzenia,

- folia miedziana jest wybierana jako większa, chroniąc w ten sposób narzędzie prasujące przed wyciekającą żywicą i eliminując potrzebę stosowania dodatkowych folii ochronnych.

Wytworzona w ten sposób technologia BASIC stanowi podstawę wielu zastosowań, na przykład w branży medycznej, motoryzacji, lotnictwie czy przemyśle.

Jeśli BASIC nie wystarcza?

Rozwiązania dostosowane do indywidualnych potrzeb zapewniają poprawę w zakresie miniaturyzacji, wydajności, niezawodności, bezpieczeństwa funkcjonalnego i kosztów systemu. Bazując na tych samych procesach, dodając kolejne procesy i dostosowując materiały i stosy, oferujemy:

- płytki drukowane HDI z mikroprzelotkami wycinanymi laserowo,
 - płytki drukowane z potencjometrami, rezystorami, przełącznikami i powierzchniami grzejnymi z zastosowaniem druku polimerowego,
 - sztywno-elastyczne płytki drukowane ze zintegrowanym elastycznym okablowaniem
 - płytki drukowane z wbudowanymi komponentami,
 - płytki drukowane o zwiększonej wymianie ciepła i naklejonym radiatorom,
 - a nawet rozciągliwe obwody drukowane.
- Zainteresowany? Skontaktuj się z:

Tomasz Renkiel, +48 785 085 700
Regional Sales Manager Poland
Wurth Elektronik CBT
International GmbH

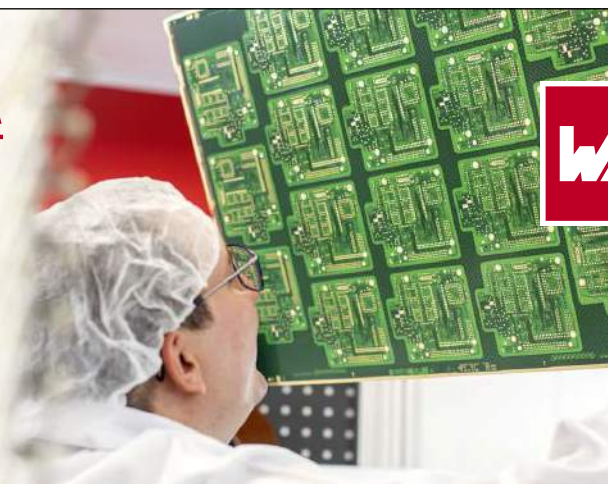


- Podstawowe etapy produkcji płytki drukowanej w filmie: www.we-online.com/video_basic_en
- Nasze standardowe stackupy BASIC: www.we-online.com/basic-stackups
- Nasze podstawowe zasady projektowania: www.we-online.com/designrulesbasic_en
- Zamów fizyczną próbkę PCB, aby zrozumieć działanie technologii BASIC: www.we-online.com/basic/sample

REKLAMA

KOMPETENCJA W PRODUKCJI PCB OD 1971

Od standardowych laminatów wielowarstwowych do technologii high-end!



**WÜRTH
ELEKTRONIK**
MORE THAN
YOU EXPECT

www.we-online.com



Design your PCB

We provide multi-lingual support
and documentation and
fast & easy cost calculation



Order your PCB

Our powerful free tools allow
virtual manufacturing
for DFM before ordering



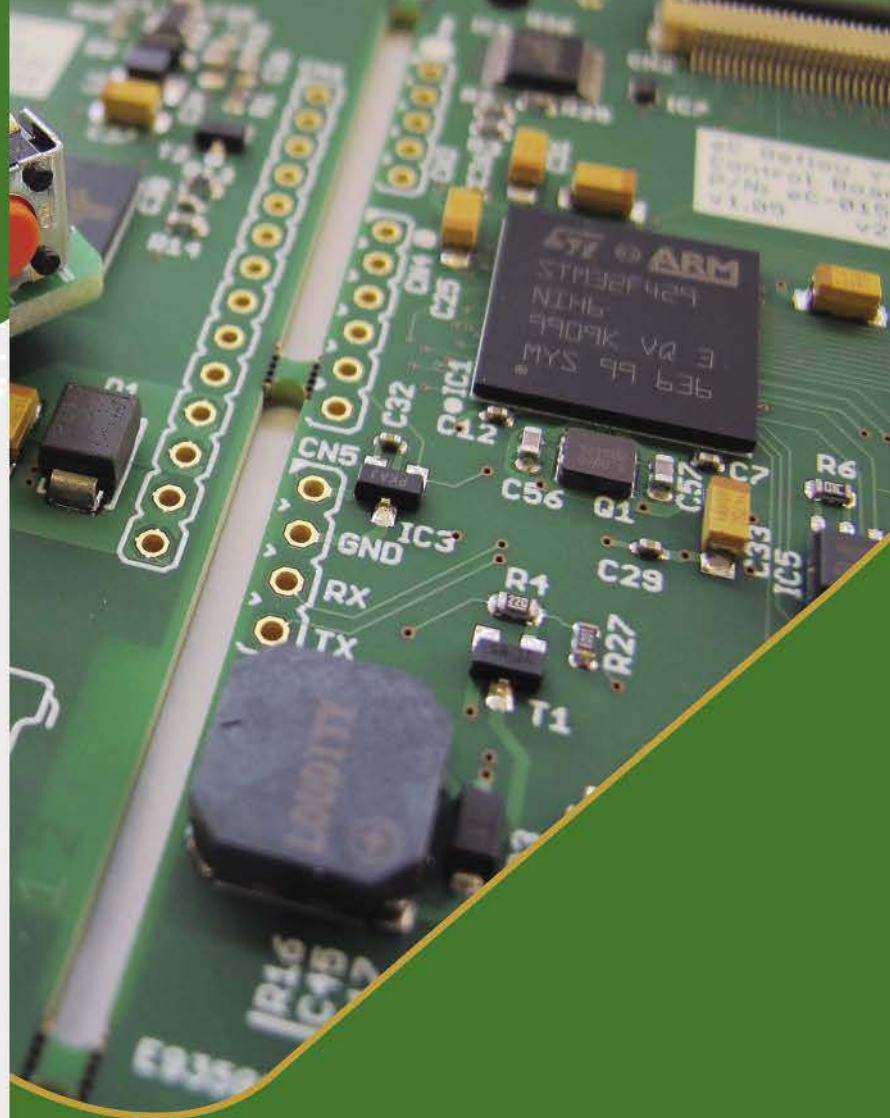
Manufacturing

Your PCB
Manufactured and Assembled
in our European factories



Right First Time

We deliver your assembled PCBs
after 6 working days

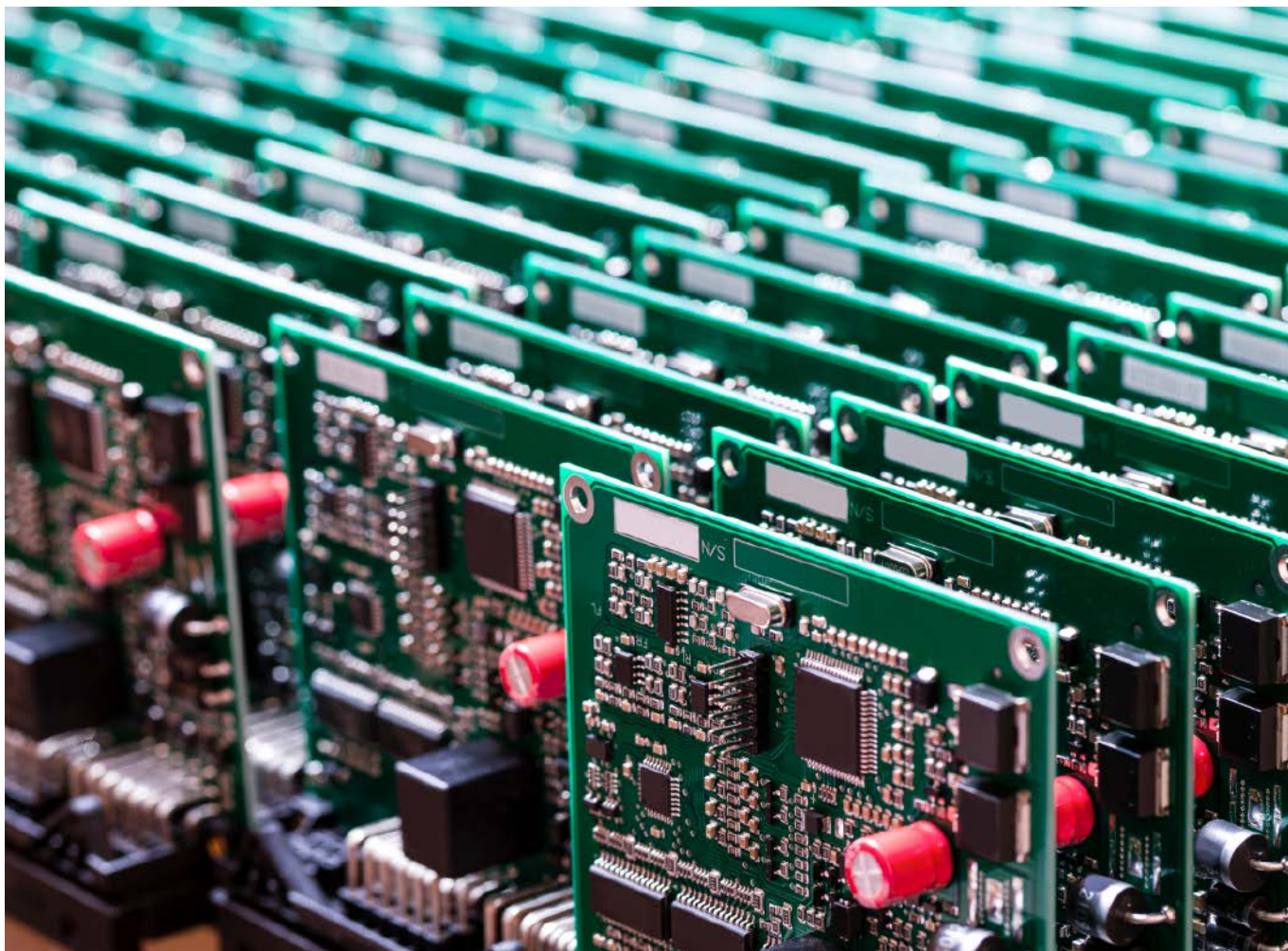


Fast & Easy ➤

PCB Prototypes & Small Series

Scan below to find news & information
about Eurocircuits and the PCB industry!





Płytki PCB zoptymalizowane do produkcji

Twój projekt PCB jest już bliski ukończenia, planujesz na początek kilka prototypów, później seria pilotowa, a potem może masowa produkcja? Zanim do tego dojdzie, warto poznać kilka istotnych zaleceń projektowych, aby uniknąć najczęstszych błędów utrudniających montaż płytki i niepotrzebnie zwiększających koszty.

Opisane zalecenia dotyczą w szczególności niedużych partii produkcyjnych, przy dużych seriach wiele aspektów wygląda inaczej, ale to nie nasza specjalność.

1. Unikaj długich, wystających na zewnątrz pól lutowniczych pod komponentami 2-wyprowadzeniowymi

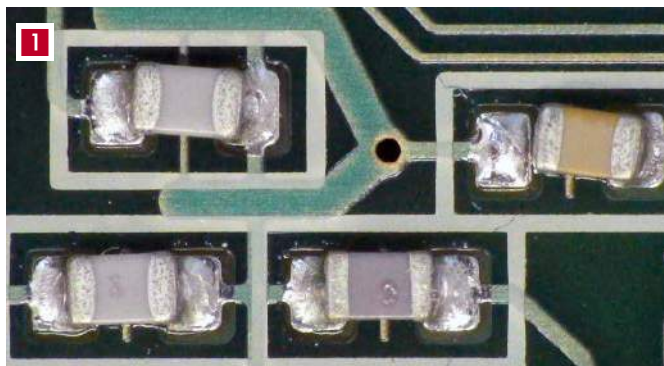
W wielu programach do projektowania PCB dostępne są specjalne biblioteki footprintów dla komponentów RLC opisane jako *hand solder*. Taki footprint ułatwia montaż ręczny, ale jest źródłem problemów przy lutowaniu rozplwowym. Podczas lutowania komponent może zostać ściągnięty przez rozgrzane spoiwo na środek jednego pola, a drugie wyprowadzenie zostanie nieprzylutowane (fotografia 1).

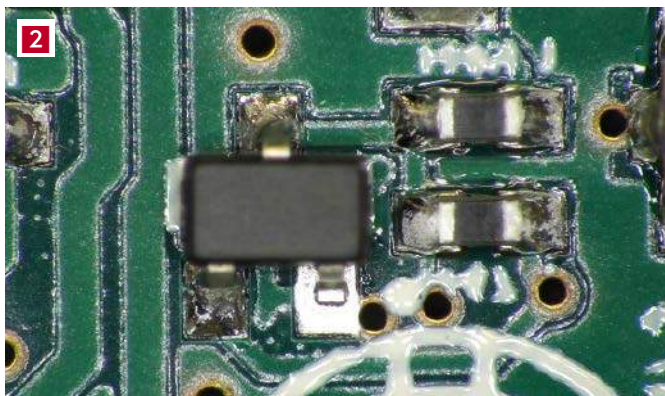
Więcej informacji:

Zakład Elektroniczny SECURUS

62-030 Luboń, ul. 11 Listopada 73, tel. 61-8331545
marekp@securus.com.pl, www.securus.com.pl

Montaż kontraktowy SMD i THT. 38 lat doświadczenia w produkcji elektroniki. Specjalność: małe serie i prototypy. Zapewniamy podstawowe komponenty, szablony do pasty. Krótkie terminy realizacji.





2. Unikaj przelotek w polach (padach) pod wyprowadzeniami komponentów

Czasem wydaje się to wygodne, gdy jest mało miejsca na płytce, ale jest to źródło potencjalnych problemów. Podczas lutowania rozgrzane spoiwo z pola wpływa do otworu przelotki i często pod wyprowadzeniem komponentu zostaje tylko jego minimalna ilość. Takie połączenie łatwo się uszkadza i może pęknąć na skutek naprężeń w płytce (**fotografia 2**).

3. Staraj się zmieścić wszystkie komponenty na jednej stronie płytki

Jeśli masz wystarczająco dużo miejsca na płytce, to jednostronne rozmieszczenie elementów da wiele korzyści. Umieszczenie kilku komponentów po drugiej stronie może ułatwić projektowanie, ale wówczas cały proces montażu – nakładanie pasty, montaż komponentów i lutowanie musi przebiegać dwa razy. Dodatkowo wydłuża się cały proces przygotowawczy, który przy małych seriach potrafi zająć więcej czasu niż sam montaż.

4. Przemyśl sposób lutowania komponentów THT

Najwygodniejszym i najszybszym rozwiązaniem jest użycie klasycznej fali lutowniczej. Wówczas projekt wymaga umieszczenia na jednej stronie płytki komponentów SMD i THT, ale nie zawsze jest to możliwe. Jeśli komponenty SMD znajdują się po stronie lutowania wyprowadzeń THT, to musimy zaplanować jeden z dodatkowych procesów:

a. Klejenie komponentów SMD

Klej nakładany jest specjalnym dyspenserem. Po tym zabiegu lutowanie komponentów SMD i THT może odbywać się jednocześnie na fali lutowniczej. Niestety miniaturyzacja komponentów THT spowodowała odejście od tej technologii.

b. Lutowanie ręczne

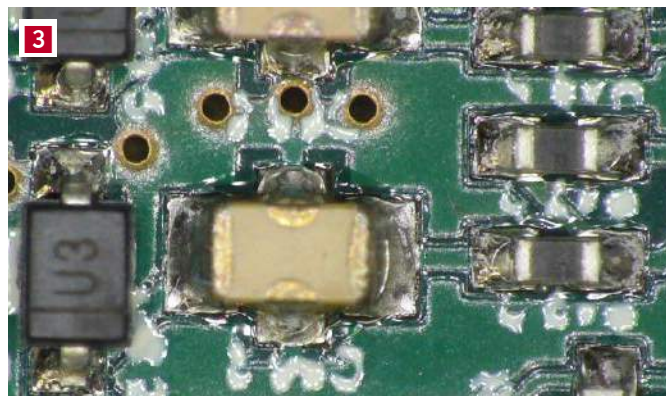
To najprostsze, ale najbardziej czasochłonne rozwiązanie. Proszę pamiętać, że na przestrzeni ostatnich 20 lat radykalnie zmieniły się proporcje między kosztami komponentów i materiałów a kosztami pracy ludzkiej. Należy w miarę możliwości minimalizować liczbę komponentów THT, chociaż do całkowitego ich wyeliminowania droga wydaje się daleka.

c. Zastosowanie selektywnej fali lutowniczej

To rozwiązanie wydaje się bardzo atrakcyjne, ale też jest kosztowne i czasochłonne. Lutowanie odbywa się w atmosferze azotu i pozwala uzyskać dobrej jakości i powtarzalne połączenia. Jednak koszt filtracji azotu z powietrza jest dość znaczny, poza tym czas przygotowania programu i testy jego działania nie kwalifikują tego rozwiązania do małych serii płytek (poniżej 100 sztuk).

d. Robot lutowniczy

To stosunkowo nowe i coraz częściej stosowane rozwiązanie też niezbyt dobrze nadaje się do małych serii. Płytki muszą być zainstalowane komponentami THT w dół, robot bowiem lutuje od góry płytki, a komponenty THT muszą być zabezpieczone przed wypadnięciem. Odbywa się to za pomocą specjalnie przygotowywanych fixtur, które dociskają elementy od spodu płytki. Przygotowanie takiej fixtury może być nieopłacalne dla małych serii.



5. Zwracaj uwagę na średnice otworów dla komponentów THT

Gotowe biblioteki komponentów zawarte w programach do edycji PCB nie zawsze zawierają właściwe średnice. W szczególności dotyczy to listew kołkowych, tzw. goldpinów. W przypadku najczęściej używanych, w rozstawie 2,54 mm otwory powinny mieć średnicę 0,9 mm. Mniejsza średnica nie pozwoli na wciśnięcie kołków, ale najczęściej spotyka się otwory o zbyt dużych średnicach. Skutkuje to problemami w lutowaniu zarówno ręcznym, jak i metodą na fali. Często stosowane odcinki o 2 lub 3 szpilkach wtykane w zbyt duże otwory mają tendencję do wypadania lub przechylania podczas transportu i lutowania. Oprócz wrażliwości na wstrząsy płynna cyna podczas lutowania na fali potrafi wypchnąć krótkie odcinki listew, zalewając przy tym otwory.

6. Stosowanie pojedynczych płytek bez panelizacji

Przy prototypowych seriach płytek można dopuścić takie rozwiązanie, jednak należy zwracać uwagę na komponenty leżące blisko krawędzi płytki. Jeśli ich odległość do krawędzi jest mniejsza niż 3 mm, istnieje ryzyko przesuwania komponentów podczas transportu. Rozwiązaniem jest zamawianie płytek z ramką technologiczną o szerokości najczęściej 5 mm.

7. Warstwa opisowa – zastanów się, czy jest potrzebna?

Przy dużym stopniu upakowania elementów na płytce często brakuje miejsca na jakiegokolwiek opisy lub obrysy komponentów. Rozdzielczość technologii nakładania opisów jest stosunkowo niewielka i często z planowanych napisów pozostają tylko bezkształtne plamy (**fotografia 3**).

8. Duży rozmiar panelu (formatki) przyspiesza montaż

Jednak duża formatka może też generować większe koszty przy małej serii. Cena szablonu do nakładania pasty zależy od liczby wycinanych laserem otworów. Duży rozmiar panelu także sprawia problemy przy stosowaniu cienkich laminatów poniżej 1 mm. Laminat taki ma tendencję do wyginania się, co utrudnia montaż na każdym etapie. Warto skonsultować ten aspekt z firmą, której zleczysz montaż, zanim zamówisz gotowe panele.

Marek Pyżalski





Kompleksowa produkcja elektroniki z zastosowaniem wielowarstwowych PCB

Wraz z postępem technologii miniaturyzacja elektroniki stała się nieodzownym elementem przemysłu. Urządzenia stają się coraz mniejsze, bardziej wydajne i funkcjonalne. Ten dynamiczny rozwój wymusza również ewolucję w projektowaniu i produkcji PCB. W dzisiejszych czasach płytki wielowarstwowe odgrywają kluczową rolę w świecie elektroniki, gdzie tradycyjne PCB o jednej lub dwóch warstwach przewodnika często stają się niewystarczające do uzyskania pożądanego funkcjonalności.

Projekty elektroniczne, które uwzględniają aspekty montażu, takie jak wybór odpowiednich komponentów, optymalizacja układu czy dostosowanie do procesów produkcyjnych, mają znaczący wpływ na finalny sukces rynkowy. Dlatego też zrozumienie tych procesów jest kluczowe dla inżynierów i projektantów, którzy dążą do opracowania konkurencyjnych produktów o najwyższej jakości, funkcjonalności i niezawodności.

Dzięki odpowiednim krokom podejmowanym na każdym etapie, firmy takie jak EAE Elektronik osiągają najwyższe standardy jakości i dostarczają produkty, które spełniają oczekiwania klientów.



Więcej informacji:

EAE Elektronik Spółka z o.o.

38-500 Sanok, ul. Przemysła 24d, www.eae-elektronik.pl



Głównym obszarem naszej działalności są kompleksowe usługi montażu kontraktowego elektroniki (EMS). W skład naszej obecnej infrastruktury wchodzi trzy pełne linie montażowe SMT, dwa zaawansowane automaty do lutowania selektywnego THT, a także jeden automat do lutowania na fali laminarnej. Dodatkowo firma dysponuje dwiema pełnymi liniami do selektywnego nakładania powłok lakierniczych, wyposażonymi w piece do utwardzania zarówno termicznego, jak i UV.

W artykule opiszemy szczegółowo kolejne etapy, które kształtują produkcję elektroniki, a także zaprezentujemy przykłady, które pomogą lepiej zrozumieć procesy występujące w firmie EAE Elektronik.

Procesy w produkcji elektroniki

Warto zaznaczyć, że zastosowanie opisanych dalej procesów może być uzależnione zarówno od złożoności projektu, wymagań klienta, jak

i warunków późniejszej eksploatacji. Skupimy się tylko na kluczowych procesach montażu elektroniki, które występują najczęściej i stanowią integralną część produkcji elektroniki z zastosowaniem płytek wielowarstwowych.

Za początek montażu elektroniki można uznać już samo zamówienie, gdzie klient przesyła pliki z dokumentacją produkcyjną. Szeroko akceptowanym standardem w branży produkcji PCB jest format Gerber wraz z plikami wierceń. Pliki w tym formacie określają niezbędne aspekty fizyczne PCB, takie jak geometrie poszczególnych warstw sygnałowych, warstwy opisowe i izolacyjne. Do wykonania montażu niezbędne jest przesłanie jeszcze dwóch dodatkowych plików: BOM (*Bill Of Materials*) – zawierającego listę komponentów i P&P (*Pick and Place*) zawierającego dane lokalizacji i orientacji komponentów.

Pliki wejściowe są sprawdzane przez zespół specjalistów pod kątem dostępności podzespołów i procesów niezbędnych do wykonania montażu. Zespół inżynierów szuka również komponentów wymagających unikalnych technik montażowych lub specyficznych profili termicznych. Po zatwierdzeniu wykonalności i akceptacji oferty przygotowywane są pliki i programy dla poszczególnych maszyn produkcyjnych. Dla optymalizacji produkcji pojedyncze płytki są często łączone w większe panele, do których dodawane są marginesy w celu ułatwienia manipulacji i transportu.

Na tym etapie tworzone są również szablony do nakładania pasty lutowniczej w procesie SMT. Szablony mają formę cienkiego arkusza ze stali nierdzewnej z otworami w miejscach wyprowadzeń komponentów.

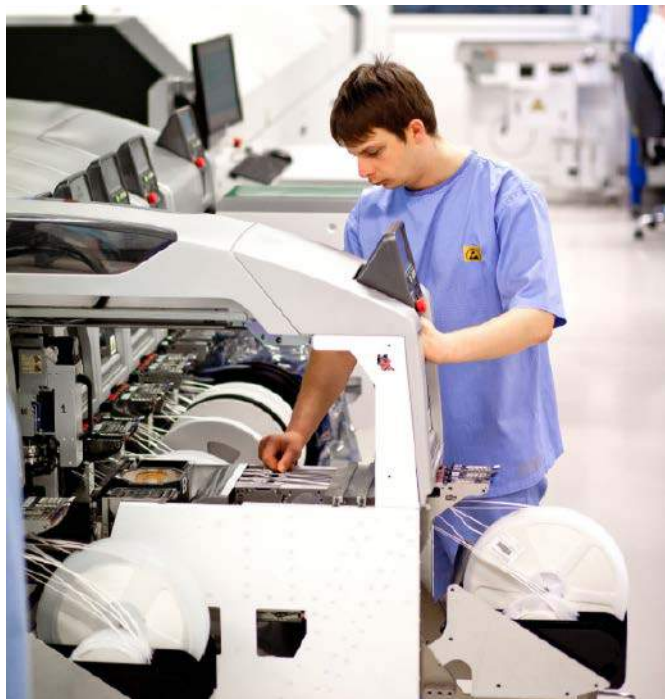
Produkcja samych PCB jest niezwykle złożonym procesem, szczególnie w przypadku płyt wielowarstwowych. Wymaga precyzyjnych procesów laminacji, w których różne warstwy PCB są łączone ze sobą pod wysokim ciśnieniem i temperaturą. Dodatkowo procesy chemiczne, takie jak wytrawianie ścieżek czy depozycja miedzi, są niezbędne do tworzenia pożądanych wzorów przewodzących na płycie. Z tego powodu produkcją PCB zazwyczaj zajmują się specjalistyczne firmy, które mają odpowiednie doświadczenie, sprzęt oraz wiedzę.

Rozpoczęcie montażu

Uruchomienie montażu elektroniki na linii produkcyjnej zaczyna się od uzbrojenia podajników maszyn Pick and Place w komponenty. Na tym etapie weryfikowane są również profile termiczne pieców i uzbrajane są maszyny do nakładania pasty lutowniczej.



Fotografia 1. Nakładanie pasty lutowniczej w procesie montażu SMT



Fotografia 2. Uzbrajanie automatów Pick and Place

Nakładanie pasty za pomocą sitodruku jest preferowaną metodą dla produkcji seryjnej z uwagi na dużą precyzję, powtarzalność i wydajność.

Płytki PCB są wyjmowane z hermetycznych opakowań zapobiegających dostawianiu się wilgoci i są umieszczane w loaderze, z którego będą automatycznie podejmowane przez maszynę do nakładania pasty. Automat do sitodruku precyzyjnie pozycjonuje szablon za pomocą tzw. fiducjali. Są to specjalne punkty rozpoznawane optycznie przez maszynę, a umieszczane na PCB podczas projektowania lub przygotowywania panelu. W procesie sitodruku maszyna przeciska pastę lutowniczą za pomocą rakli przez mozaikę otworów szablonu – **fotografia 1**. Nowoczesne automaty takie jak MPM Momentum II używane w EAE Elektronik pozwalają na kontrolę wielu parametrów tego procesu, takich jak siła nacisku czy temperatura pasty, dla zapewnienia wysokiej precyzji przy możliwie najkrótszym czasie cyklu pracy.

Bezpośrednio po nałożeniu pasty następuje jej inspekcja w automacie SPI (*Solder Paste Inspection*). Nasz automat Parmi Sigma X wykonuje laserowy skan 3D, mierząc ilość i pozycję nałożonej pasty. Dalszy proces montażu jest wykonywany tylko na panelach z idealnie naniesioną pastą.

Płytki, które pozytywnie przeszły test SPI, są transportowane do automatów układających komponenty montowane powierzchniowo. W firmie EAE Elektronik używamy między innymi automatów Pick and Place Fuji NXT III. Większość komponentów SMD jest podejmowana z rolowanych taśm za pomocą głowic z końcówkami podciśnieniowymi. Komponenty mogą być również pobierane z tacek i szyn – **fotografia 2**. Dla komponentów o małej masie, jak rezystory i kondensatory, stosuje się głowice pozwalające na jednoczesne pobieranie i przenoszenie nawet do 24 elementów. Oszczędzamy w ten sposób na czasie przejazdu od podajnika do PCB, co w optymalnych warunkach pozwala na osiągnięcie wydajności do kilkudziesięciu tysięcy komponentów na godzinę (cph).

Systemy optyczne automatów P&P dbają o prawidłowe pozycjonowanie i rotację komponentów, porównując je podczas każdego pobrania z wzorcami zapisanymi w pamięci.

Proces lutowania

Następnie płytki są transportowane w kierunku pieca do lutowania rozplwowego. Przechodząc kolejne etapy, panele PCB w pierwszej kolejności trafiają do strefy rozgrzewania wstępnego i stabilizacji

PREZENTACJE

temperatury. Jej celem jest usunięcie resztek wilgoci i równomierne nagrzanie płytki oraz komponentów. Następnie płytki trafiają do strefy rozplwy o wyższej temperaturze, gdzie następuje stopienie pasty lutowniczej i tworzą się właściwe połączenia elektryczne. Ostatnim etapem jest chłodzenie, które przeprowadzane w kontrolowany sposób pozwala na uzyskanie odpowiedniego rozkładu cyny w spoinie. Maksymalne temperatury osiągane w strefie rozplwy zależą od rodzaju użytej pasty i wahają się zazwyczaj od ok. 180°C dla past niskotemperaturowych, np. bizmutowych, do ok. 260°C dla past bezołowiowych.

Automatyczna inspekcja optyczna

Po lutowaniu SMT panele PCBA (*Printed Circuit Board Assembly*) zawierające komponenty są weryfikowane w procesie automatycznej inspekcji optycznej (AOI). Etap ten pozwala na wychwytywanie błędów związanych z brakiem lub nieprawidłowym ułożeniem komponentów oraz wad lutowania. Odbywa się to przez wykonanie serii zdjęć wysokiej rozdzielczości i automatyczne porównanie ich z prawidłowymi wzorcami.

Nowoczesne maszyny stosowane w EAE Elektronik łączą funkcje 2D oraz 3D. Szybkie laserowe skanowanie płytki pozwala rozszerzyć funkcje rozpoznawania obrazu o badanie profilu wysokościowego. Przekłada się to na dużą dokładność i wykrywalność błędów, które mogłyby zostać niezauważone przez standardowe techniki rozpoznawania obrazu.

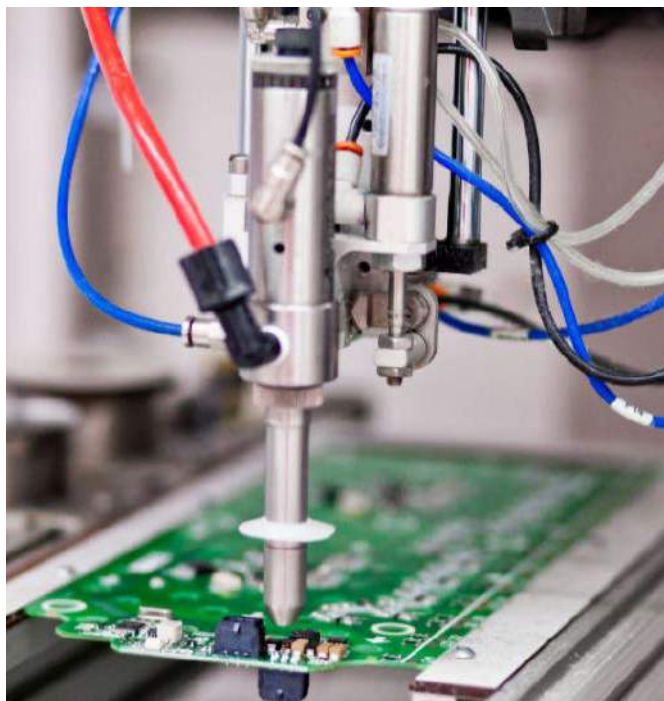
Montaż elementów przewlekanych

W większości przypadków zespoły PCBA zawierają również komponenty montowane w technice przewlekanej THT. Montaż tą techniką charakteryzuje się przede wszystkim wyższą stabilnością mechaniczną. Komponenty są mocno osadzone w otworach PCB i lutowane od spodu, co sprawia, że są bardziej odporne na wibracje oraz wstrząsy – **fotografia 3**.

Stosowanie płytek wielowarstwowych często wiąże się z montażem dwustronnym, gdzie komponenty są umieszczane na obu stronach PCB. Stosowanie maszyn do lutowania selektywnego THT jest wówczas najczęściej wybieraną metodą. Automat do lutowania selektywnego ma tygiel wyposażony w głowicę z dyszą, pozwalającą na punktowe lutowanie. Aby zapewnić odpowiednią zwilżalność i jakość połączeń lutowniczych, proces standardowo przeprowadzany jest w atmosferze osłony azotowej.



Fotografia 3. Układanie komponentów THT



Fotografia 4. Automatyczna linia lakiernicza

Jeśli proces montażu tego wymaga, po zakończeniu lutowania może zostać wykonane mycie płytek PCBA. Jego celem jest usunięcie pozostałości topnika oraz innych zanieczyszczeń powstałych podczas procesu montażu i produkcji. Mycie może również być elementem procesu przygotowania PCBA do nałożenia powłoki ochronnej. Należy jednak zwrócić uwagę, że istnieją komponenty takie jak mikrofony czy elementy elektromechaniczne, które mogą być wrażliwe na wilgoć i chemikalia używane podczas procesu mycia, co wyklucza niektóre projekty z tego procesu.

Dodatkowe powłoki

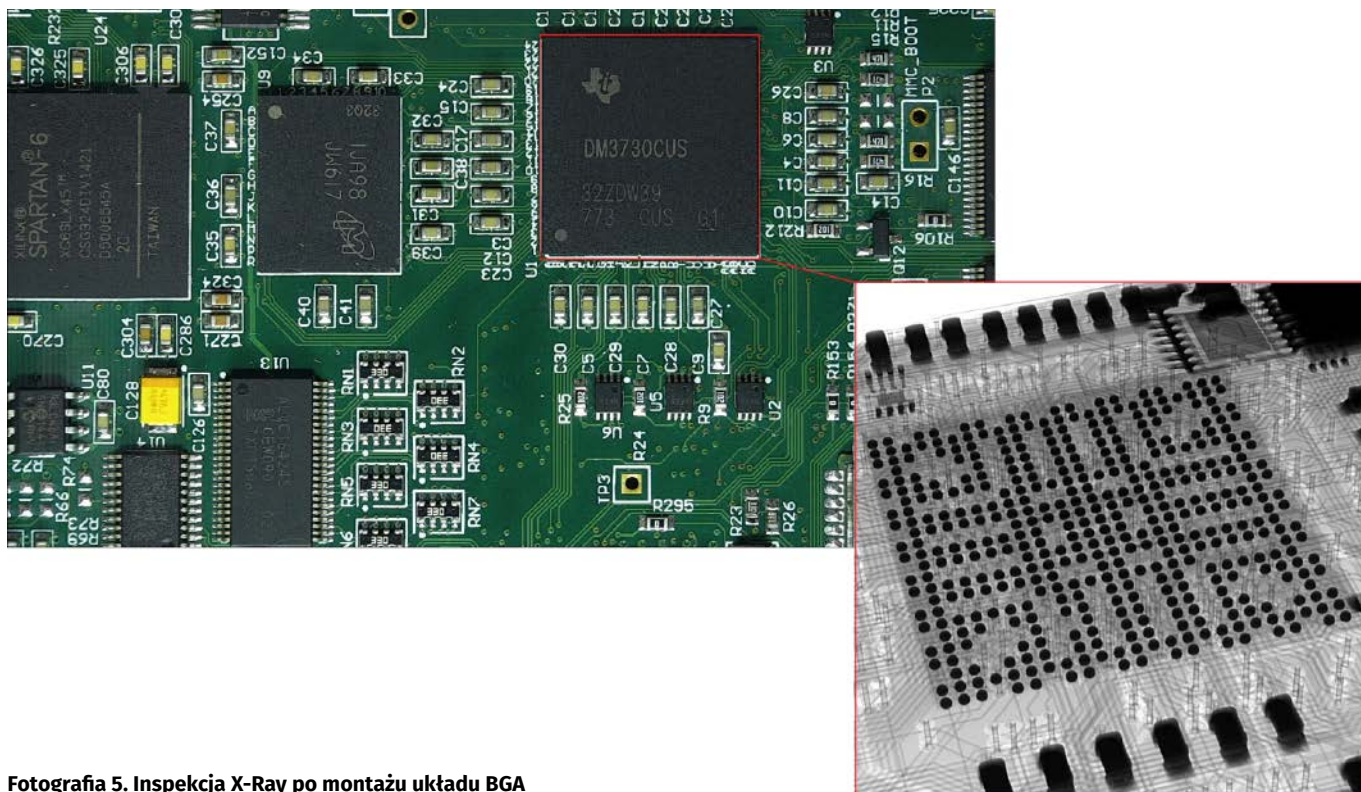
Kolejnym opcjonalnym procesem jest nakładanie powłok konforemnych. Lakier chroni przed wilgocią i zmniejsza ryzyko rozrostu wąsów cynowych (*whiskers*), znacząco wydłużając żywotność produktu. W procesie nakładania powłok ochronnych istotne są zarówno zastosowana metoda jak i rodzaj użytego lakieru. W firmie EAE Elektronik oferujemy standardowo trzy rodzaje lakierów. Najpopularniejszym są lakiery akrylowe, których utwardzanie polega na odparowaniu rozpuszczalników. Dzięki temu proces ten jest stosunkowo szybki i może zostać dodatkowo przyspieszony przez wygrzewanie płytek w piecu.

W przypadku zastosowań, w których wymagana jest podwyższona odporność na czynniki chemiczne, proponujemy lakiery poliuretanowe, utwardzane światłem UV. Natomiast tam, gdzie konieczna jest odporność na wysokie temperatury, wybierane są lakiery silikonowe. Selektywne nakładanie powłok konforemnych umożliwia zabezpieczenie tylko tych części, które tego wymagają, bez konieczności pokrywania całej powierzchni PCBA – **fotografia 4**.

Procesy końcowe

W końcowych etapach produkcji zwykle następuje separacja paneli na pojedyncze płytki. W trakcie separacji pojawia się ryzyko narażeń mechanicznych, mogących wpłynąć na integralność płytek wielowarstwowych oraz zamontowanych komponentów. Dlatego w firmie EAE Elektronik stosujemy specjalistyczne urządzenia do separacji oraz narzędzia do wyłamywania mostków, co umożliwia precyzyjne oddzielenie PCBA bez ryzyka uszkodzeń.

Warto również podkreślić, że w trakcie produkcji na PCBA mogą być wykonywane operacje programowania i różnego rodzaju testy jakości. Obejmują one m.in. testy ICT (*In-Circuit Test*), weryfikujące



Fotografia 5. Inspekcja X-Ray po montażu układu BGA

poprawność połączeń i funkcjonalność komponentów, oraz testy FCT (*Functional Test*), sprawdzające funkcjonalność urządzenia. Testy EOL (*End-of-Line*) stanowią natomiast ostatni etap w procesie produkcji, gdzie dokładnie weryfikuje się gotowe produkty pod kątem poprawności ich działania i spełnienia wymagań klienta.

W celu zapewnienia najwyższej jakości usług firma EAE Elektronik zainwestowała w aparat General Electric Microme X do pomiarów i inspekcji rentgenowskiej (X-ray), który jest stosowany na różnych etapach produkcji. Jest niezwykle przydatny podczas walidacji procesu montażu SMT. Daje możliwość obserwacji ukrytych połączeń pod obudową, radiatorem lub warstwą lakieru ochronnego.

W kontekście płytek wielowarstwowych, inspekcja X-ray stanowi również jedyne narzędzie pozwalające w sposób nieniszczący na zajrzenie do wnętrza PCB. Dzięki temu możliwa jest weryfikacja defektów, takich jak pęknięte ścieżki, uszkodzone przelotki, czy stopień wypełnienia lutownikiem w montażu THT. Urządzenie ma również

funkcję generowania skanów 3D CT (Computed Tomography), pozwalających na analizę obiektów z dowolnej perspektywy i przybliżenia – **fotografia 5**.

Ciągła kontrola procesów produkcyjnych i dogłębna analiza przyczyn powstawania usterek odgrywają kluczową rolę w zapewnieniu najwyższej jakości produktów. Dlatego też firma EAE Elektronik dysponuje własnym zaawansowanym laboratorium badawczym, które stanowi istotne uzupełnienie oferty w zakresie usług EMS. W laboratorium przeprowadzamy szeroki zakres badań, takich jak analiza kompatybilności elektromagnetycznej, badania funkcjonalne, trwałościowe, oraz testy bezpieczeństwa – **fotografia 6**. Dzięki bogatemu zapleczu i wykwalifikowanej kadrze jesteśmy w stanie wykrywać potencjalne nieprawidłowości na wczesnym etapie projektowania i produkcji.

Nasze badania stanowią również duże wsparcie dla klientów. Nierzadko udaje się wskazać konkretne zmiany projektowe, które przyczyniają się do poprawy jakości i niezawodności produktów zleconych do montażu.



Fotografia 6. Test odporności na wyładowania elektrostatyczne (ESD)

Podsumowanie

Inżynierowie i projektanci muszą uwzględniać procesy produkcyjne już na etapie projektowania, aby zapewnić powtarzalność produkcji i osiągnąć wysoką jakość. Należy myśleć perspektywnie, uwzględniając nie tylko jeden prototyp, lecz tysiące, a nawet miliony egzemplarzy. W praktyce projektowanie płytek wielowarstwowych wymaga zastosowania podejścia optymalizacji dla produkcji (DFM). Zrozumienie i przeciwdziałanie niekorzystnym zjawiskom podczas produkcji odgrywają istotną rolę w zwiększaniu konkurencyjności, oraz osiągnięciu sukcesu na rynku.

M. Książek
R&D EAE Elektronik

Laminat FR-4

– używany przez wielu, znany przez nielicznych

Najpopularniejszy rodzaj laminatu do płytek PCB to FR-4. Jest nieprzewodzącym materiałem wykonanym z włókien szklanych zatopionych w żywicy epoksydowej, który ma właściwości najbardziej odpowiednie dla szerokiego zakresu zastosowań PCB. Choć jest zdecydowanie najpopularniejszym rozwiązaniem, to daleko mu do prostego i jednolitego materiału.

Laminat FR-4 to w rzeczywistości kategoria materiałów zdefiniowana przez NEMA od lat 60. XX wieku. Odpowiednia norma określa wymagania eksploatacyjne FR-4, które znacząco ewoluowały na przestrzeni lat. Wiele projektantów płytek PCB nawet nie ma świadomości tego, że ten pozornie podstawowy element jest przedmiotem stałych badań i rozwoju. Jest to konieczne, aby sprostać rosnącym wymaganiom przemysłu elektronicznego, zwłaszcza w zakresie optymalizacji surowców, a w szczególności w zakresie redukcji zmienności parametrów tego produktu.

Najważniejsze właściwości, których oczekujemy od laminatu FR-4, to:

- zawiera spójne, jednolite i przewidywalne właściwości nawet po kilku cyklach ogrzewania i chłodzenia, podczas których materiał mięknie, płynie, a następnie ponownie twardnieje;
- miękki na tyle, aby wypełnić maleńkie szczeliny w miedzi i jednocześnie odpowiednio wytrzymały;
- tworzy mocne i trwałe wiązania, podobnie jak klej, pomiędzy warstwami technologicznymi (tj. pomiędzy laminatami, miedzią i maską lutowniczą);
- dostępny w postaci sztywnych rdzeni i bardzo cienkich arkuszy, o różnych właściwościach mechanicznych i długim okresie trwałości. To naprawdę sporo, jeśli chodzi o materiał!

Aby wyprodukować niezawodne płytki PCB, jakoś surowców musi być na wysokim poziomie, a proces produkcyjny musi być specjalnie dostosowany do użytego laminatu. Ponadto należy wybrać materiał o odpowiednich właściwościach, aby uwzględnić przyszłe zastosowanie płytki – np. jeśli będzie przechodziła wiele cykli ogrzewania i chłodzenia lub będzie działała w środowiskach o wysokiej temperaturze. W artykule opisujemy różne aspekty jakościowe laminatu FR-4, aby projektanci PCB byli świadomi potencjalnych punktów krytycznych i mogli omówić je z producentami w celu wybrania najlepszych materiałów do swojego zastosowania.

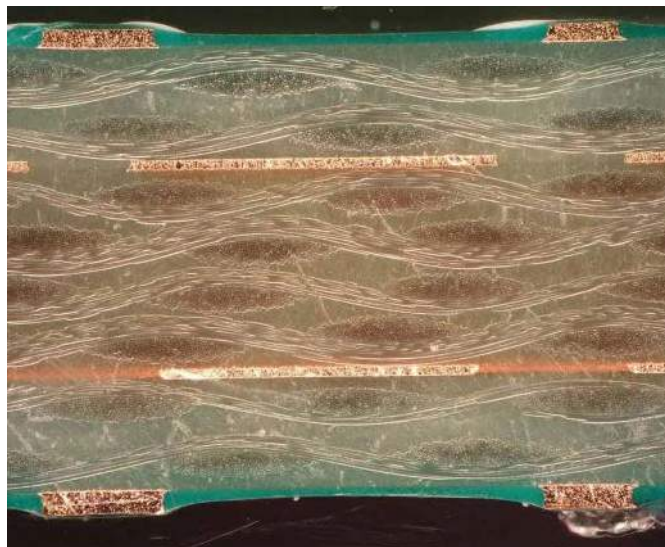
Patrząc na przekrój czterowarstwowej płyty (**rysunek tytułowy**), można odkryć wiele szczegółów. Widzimy pojedyncze włókna szklane tworzące wiązkę i to, że wiązki te tworzą splot, który jest zatopiony w żywicy. Podczas procesu tworzenia warstw płytki PCB dodawana jest miedź, aby utworzyć obwód elektryczny.

Jak używa się laminatu FR-4?

FR-4 występuje w dwóch podstawowych postaciach:

- **Rdzeń** – utwardzony laminat z jednym lub dwoma arkuszami miedzi na powierzchni. Jest to podstawa do wykonania dwuwarstwowej płytki PCB.
- **Prepreg** – skrót od preimpregnat – częściowo utwardzone arkusze laminatu z żywicy epoksydowej, utwardzanie i wiązanie następuje pod wpływem ciepła podczas prasowania.

Rdzenie i prepregi są dostępne w wielu grubościach, co pozwala producentom na tworzenie nieskończenie wielu warstw. Eurocircuits oferuje ponad 975 predefiniowanych warstw!



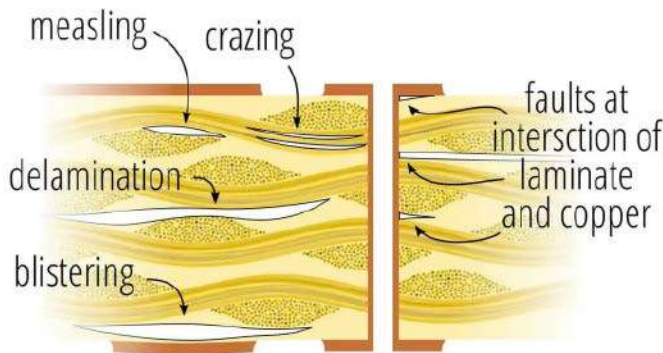
Oto bardzo uproszczony opis tworzenia układów PCB. W przypadku płytek dwuwarstwowych producenci korzystają z gotowych rdzeni, które poddawane są procesom chemicznym i mechanicznym w celu uzyskania gotowej płytki drukowanej. W przypadku czterech lub więcej warstw płyty są tworzone poprzez ułożenie wytrawionych (i nawierconych w razie potrzeby) rdzeni z arkuszami prepregu (zwykle dwoma) pomiędzy sąsiadującymi ze sobą warstwami miedzi. Zestawienie to jest podgrzewane i prasowane, co powoduje, że nieutwardzona żywica w prepregach topi się i wciska w szczeliny w miedzi i w otworach (np. w ślepych i zakopanych przelotkach). Uzyskuje się przyczepność, która nigdy nie powinna się rozerwać przy określonym użytkowaniu i w określonych warunkach. Na koniec stos przechodzi ten sam proces chemiczny i mechaniczny co płyta dwuwarstwowa. Powstały stos może teraz stanowić element konstrukcyjny do wytworzenia konstrukcji z dodatkowymi warstwami, ułożonymi razem w wyniku dalszych prasowań.

Czasami słyszy się o prepregu o niskim stopniu płynięcia lub jego braku, który, jak sama nazwa wskazuje, staje się płynny tylko w procesie łączenia warstw, ale nie wypełnia pustych przestrzeni. Ten rodzaj prepregu nie kurczy się (w pionie) ani nie rozszerza (w poziomie) podczas prasowania stosu. Przykładem zastosowania tej opcji są płyty elastyczne i półelastyczne, w przypadku których oczekujemy zachowania mechanicznych parametrów.

Co może pójść nie tak?

Oczywiście obowiązkiem producenta jest upewnienie się, że jego proces jest dostosowany do materiałów, których używa w celu wytworzenia niezawodnych płyt. Wszystko zaczyna się już od logistyki, ponieważ prepreg ma ograniczony okres przydatności do użycia, po którym ulega degradacji. Dlatego producenci muszą upewnić się, że materiały nie są przeterminowane i zachowują parametry określone w specyfikacji.

Wzory miedzi w indywidualnym projekcie mają znaczenie – więcej miedzi to mniej szczelin i grubsza warstwa; mniej miedzi oznacza więcej szczelin i cieńszą warstwę. Należy pamiętać, że żywica musi przepływać przez wszystkie te maleńkie szczeliny w miedzi – np. pomiędzy ścieżkami par różnicowych – aby płytka działała prawidłowo. Nieprawidłowe wykonanie tego etapu spowoduje problemy – przerwy



Rysunek 1. Niektóre z wad, które mogą wystąpić w FR-4

powodujące niedopasowanie impedancji, rozwarstwienie miedzi z laminatu, nierówne pokrycie otworów i wiele innych (rysunek 1).

Specyfikacja IPC-A-600 mówi nam szczegółowo, co może pójść nie tak z FR-4 i co możemy uznać za „akceptowalne” dla trzech klas PCB. Oto niektóre z tych usterek:

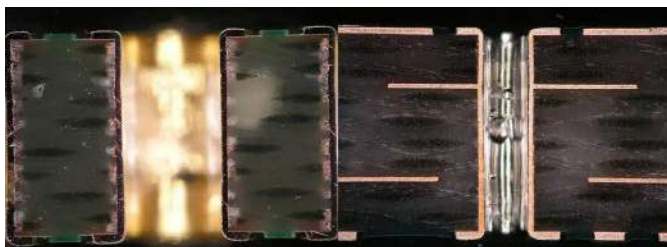
- **Measling** – oddzielenie włókien na przecięciu spłotu,
- **Crazing** – puste przestrzenie pomiędzy włóknami szklanymi w wiązce (przędzy),
- **Delamination/blistering** (delaminacja/pęcherze) – separacja w materiale bazowym, pomiędzy poszczególnymi warstwami spłotu (warstwami) lub pomiędzy materiałem bazowym, folią miedzianą lub powłoką (taką jak maska lutownicza).
- **IPC** mówi również o ekspozycji spłotu na powierzchni:
- **Weave exposure** (ekspozycja spłotu) – włókna szklane są odsłonięte z powodu braku żywicy,
- **Weave texture** (ekstura spłotu) – włókna szklane są widoczne przez cienką warstwę żywicy,
- Odsłonięte lub przerwane włókna.

Istnieje również ograniczenie dotyczące dozwolonych metalicznych lub niemetalicznych „wtrąceń obcych” uwieczonych w materiale izolacyjnym oraz ograniczenia „halo”, czyli chropowatości odsłoniętej powierzchni po wierceniu i frezowaniu laminatu. To sporo rzeczy, które mogą pójść nie tak! Omówmy teraz właściwości tych materiałów i procesy, którym są poddawane, aby zrozumieć, w jaki sposób mogą wystąpić te wady.

Współczynnik rozszerzalności cieplnej CTE

PCB przechodzi kilka cykli ogrzewania i chłodzenia. Dokładna liczba zależy od stopnia złożoności i zastosowania. W przypadku produkcji istnieje mniej więcej jeden cykl dla każdej pary warstw miedzi i dla każdej strony lutowania elementu. Następnie może wystąpić cykl podczas przeróbki lub przyłutowania modułu. I oczywiście produkt może przejść kilka lub wiele cykli, w zależności od zastosowania. Płyta musi zatem przetrwać wszystkie te cykle i resztę swojego życia.

Chociaż płytka drukowana wydaje się monolityczna, składa się z kilku materiałów, z których każdy reaguje na temperaturę w inny sposób (rysunek 2). Jedną z miar tej reakcji jest współczynnik rozszerzalności cieplnej, CTE, który mówi nam, o ile materiał rozszerza się



Rysunek 2. Przekroje przelotek w płycie dwuwarstwowej (po lewej) i wielowarstwowej (po prawej). Pęcherzyki w przelotce są zawieszono w żywicy

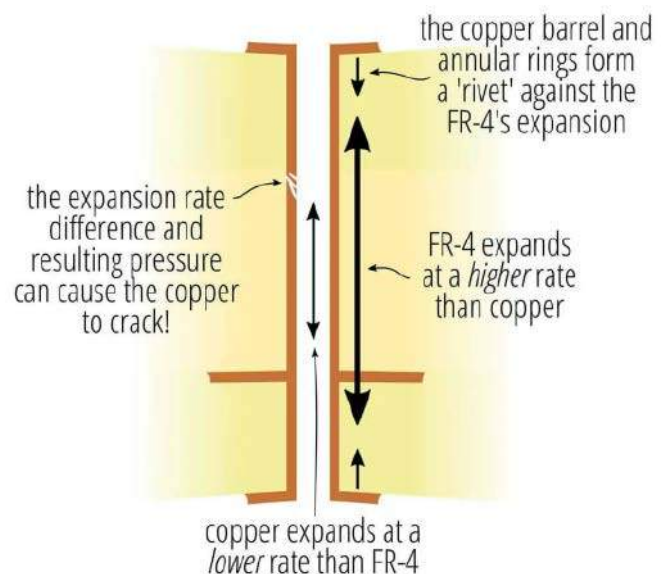
przy każdej zmianie jednostki temperatury, zwykle o 1°C. CTE jest zwykle wyrażany w ppm/°C – rozszerzanie się w częściach na milion jednostki długości na każdy stopień zmiany temperatury (czasami zobaczysz także całkowite rozszerzenie wyrażone jako procent w pewnym zakresie temperatur, na przykład 3,3% dla 50...260°C).

Jednolite materiały, takie jak miedź, będą miały pojedynczy współczynnik CTE, ponieważ rozszerzają się z tą samą szybkością we wszystkich kierunkach. Materiał kompozytowy, taki jak FR-4, powinien mieć zdefiniowane współczynniki CTE dla trzech osi X, Y i Z (odpowiednio CTE_x, CTE_y, CTE_z). Jeśli wzór spłotu w poziomych osiach X i Y jest taki sam, CTE_x jest taki sam jak CTE_y, w przeciwnym razie będą nieco inne. Jednak oś Z różni się nieco od poziomych współczynników CTE, ponieważ włókna wzmacniające przebiegają tylko w osiach X i Y. Współczynnik CTE powinien być również określony dla temperatury zeszklenia oraz poniżej i powyżej temperatury zeszklenia, T_g.

Weźmy prosty przykład użycia wartości CTE. Jeśli dla naszego laminatu CTE_x, CTE_y określono jako 17 ppm/°C, a powierzchnia materiału wynosi 50×50 mm, spodziewamy się, że materiał będzie się rozszerzał w osiach X i Y z szybkością $50/106 \times 17 = 0,00085$ mm/°C, czyli $0,00085 \times 200 = 0,17$ mm przy 200°C.

Problemy zaczynają się pojawiać, gdy materiały o różnych współczynnikach CTE są łączone ze sobą, a następnie podgrzewane. Wyobraź sobie, że przyklejasz kartę papierową (niski współczynnik CTE) do powierzchni gumki (wysoki współczynnik CTE), a następnie ciągniesz za końce opaski... papier pęknie i rozerwie się. To „rozdarcie” może również nastąpić zarówno wewnątrz płytki drukowanej, jak i na interfejsie – połączeniach lutowanych pomiędzy płytką a komponentami.

Ponieważ w osi Z nie ma zbrojenia szklanego, współczynnik rozszerzalności żywicy jest głównym współczynnikiem rozszerzalności, który ma różne wartości poniżej i powyżej temperatury zeszklenia T_g. CTE_z może wynosić nawet 70 ppm/°C poniżej T_g, ale może wzrosnąć do ponad 250 ppm/°C powyżej T_g. Prowadzi to do bardzo dużego wzrostu rozszerzalności po przekroczeniu T_g, co prawie na pewno ma miejsce podczas lutowania. Szczególny problem budzą platerowane otwory i przelotki, które tworzą miedziane nity – pierścienie na górze i na dole, połączone tuleją w laminacie. Współczynnik CTE miedzi wynosi około 17 ppm/°C, podczas gdy współczynnik CTE_z jest znacznie wyższy (CTE_x, CTE_y mają współczynnik CTE zbliżony do współczynnika CTE miedzi). Oznacza to, że wraz ze wzrostem temperatury laminat rozszerza się bardziej niż miedź, co może spowodować pęknięcie korpusu lub oderwanie się podkładek od powierzchni PCB (rysunek 3). Ryzyko to zwiększa się w przypadku



Rysunek 3. Graficzne zobrazowanie otworu przelotowego w płycie drukowanej i pęknięcia powstałego w miedzi

powtarzających się cykli i czasu, przez jaki płytka PCB znajduje się w określonych temperaturach.

Istnieje wówczas możliwość niedopasowania współczynnika CTE pomiędzy poziomym rozszerzaniem płytki PCB (CTEx, CTey) a przylutowanymi do niej komponentami. Elementy te są wykonane z różnych materiałów, które również rozszerzają się z różną szybkością; na przykład układ BGA prawdopodobnie będzie miał płytkę drukowaną wewnątrz obudowy, a połączenie z główną płytką PCB odbywa się za pomocą małych kulek lutowniczych. Niedopasowania CTE mogą powodować ścinanie i pęknięcia w połączeniach lutowniczych, co może powodować rozłączenie lub zmniejszenie długoterminowej niezawodności produktu.

Należy pamiętać, że materiały kompozytowe mogą mieć rozbudowane specyfikacje CTE, a projektanci muszą upewnić się, że biorą pod uwagę te właściwości dla ich zastosowania.

Temperatura zeszklenia Tg

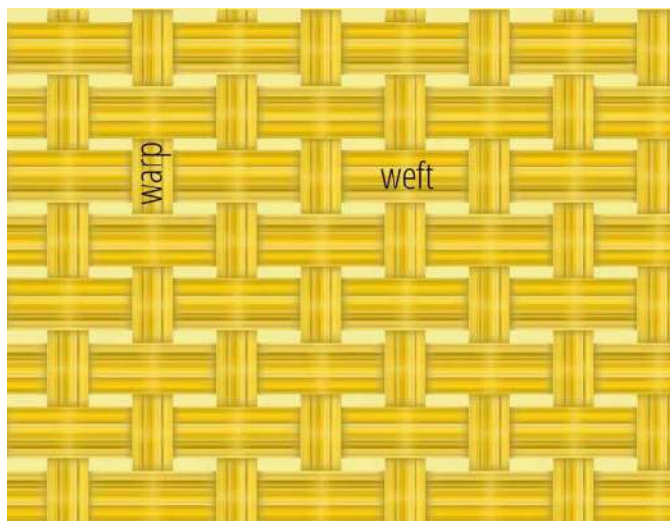
Laminaty FR-4 i inne rodzaje mają właściwość nazwaną temperaturą zeszklenia Tg. Jest to temperatura, w której polimer zmienia fazę pomiędzy stanem szklisto-kruchym a stanem lepko-gumowym. W temperaturach wyższych od Tg następuje znaczny wzrost współczynnika CTE, a co za tym idzie, zwiększenie objętości zajmowanej przez polimer. Najczęściej używany materiał FR-4 będzie miał Tg wynoszącą około 130...140°C. Inne materiały będą coraz droższe i egzotyczne w miarę wzrostu specyfikacji Tg – 180°C jest uważane za „wysoką Tg”, ale na przykład laminat z rodziny Rogers RO4000 może mieć Tg równą 280°C.

Chociaż współczynniki CTE powyżej Tg są ważne głównie dla producenta płytek PCB, sama Tg jest ważnym parametrem przy zastosowaniu płytki PCB. To tutaj projektant musi upewnić się, że temperatura otoczenia – na przykład temperatura otoczenia lub lokalna dla układu scalonego – pozostaje poniżej tej wartości z pewnym marginesem. W przeciwnym razie płytka PCB straci swoje właściwości mechaniczne i nie będzie zachowywać się zgodnie z oczekiwaniami.

Należy pamiętać, że Tg nie jest wskaźnikiem wydajności cieplnej ani wytrzymałości. Temperatura rozkładu – Td, opisana szczegółowo jest w tym przypadku bardziej odpowiednią właściwością.

Splot z włókna szklanego

Włókna szklane tworzące FR-4 nie zawsze są takie same. W rzeczywistości istnieje wiele sposobów na utkanie FR-4! Możesz zmieniać liczbę pasm włókien w wiązce, grubość każdej wiązki (można je



Rysunek 4. Każda wiązka składa się z kilku pasm włókien i jest ze sobą spleciona (jeden kierunek nazywany jest „wątkiem” (weft), a drugi „splotem” (wart)). Należy zauważyć, że odległość pomiędzy wiązkami i grubość wiązek niekoniecznie są takie same w kierunkach x i y

„spłaszczyć”), odległość między przędzami i mieć różne parametry dla osi x i y (odpowiednio „wątek” – weft i „osnowa” – warp, **rysunek 4**). Następnym parametrem jest liczba splotów na jednostkę grubości. Oznacza to, że stosunek szkła do żywicy w naszym konkretnym splocie FR-4 znacząco wpływa na wartość CTE – żywica ma wyższy CTE niż szkło, żywica jest częścią, która mięknie.

Laminat FR-4 jest anizotropowy – wykazuje różne właściwości w różnych kierunkach. Zatem pomimo tego, co moglibyśmy założyć, względna przenikalność elektryczna DK (stała dielektryczna ϵ_r , czasami określana jako „mikro-DK”) również nie jest jednolita. Jest to możliwe, skoro każde miejsce ma inny stosunek i gęstość szkła i żywicy. Może to oczywiście mieć wpływ na wynikową impedancję charakterystyczną wzdłuż ścieżek. W przypadku zastosowań o niskiej częstotliwości nie stanowi to większego problemu, ale wraz ze wzrostem częstotliwości sygnałów efekt splotu włókien (powiązany również z przekrzywieniem splotu szklanego) ulega pogorszeniu. Oczywiście problemy te zaostrzają się, gdy w FR-4 występują usterki, o których mówiliśmy wcześniej – pęcherze, rozwarstwienia i inne anomalie.

Inne właściwości

Istnieją inne specyfikacje podawane przez producentów laminatu FR-4. Temperatura rozkładu – Td, to temperatura, w której materiał traci 5% swojej masy po kontrolowanym wzroście temperatury. Jest to miara szybkości degradacji. Td jest ważnym pomiarem w przypadku montażu w procesie wymagającym wyższych temperatur, np. w przypadku materiałów bezołowiowych.

Jako miarę wydajności materiału przyjmuje się czas do rozwarstwienia. Zwykle jest określany jako minuty w trzech temperaturach 260/288/300°C – T260/T288/T300. Jak sama nazwa wskazuje, jest to czas w każdej temperaturze, po którym następuje rozwarstwienie, czyli rozdzielanie wiązań w laminacie, a w przypadku miedzi występują problemy, o których pisaliśmy wcześniej.

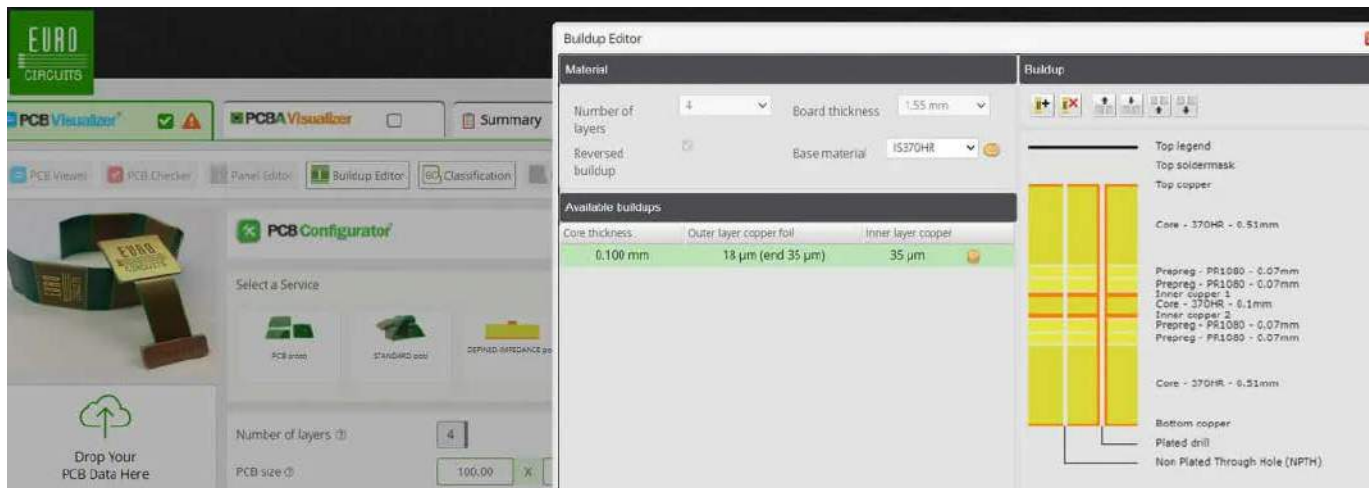
FR-4 jest wrażliwy na wilgoć, która po wchłonięciu zmienia właściwości materiału. Na przykład DK wzrasta, a Tg maleje. Wilgoć może zostać wchłonięta podczas produkcji lub w docelowej aplikacji, co może prowadzić do wielu problemów. Wilgoć jest również przyczyną uszkodzeń przewodzącego włókna anodowego – CAF, jednak te zagadnienia wykraczają poza zakres tego artykułu. Arkusz danych FR-4 powinien zawierać procentową wartość „absorpcji wilgoci” (zgodnie z definicją w IPC-TM-650 2.6.2.1), a producenci i monterzy powinni upewnić się, że ich procesy nie powodują błędów. Po wyprodukowaniu to projektanci powinni rozważyć i monitorować wpływ wilgoci na swój produkt pod względem montażu komponentów (w przypadku gołych płytek PCB), transportu i przechowywania.

Co to wszystko oznacza dla projektanta elektroniki?

Choć kusi nas, aby myśleć o FR-4 jako o jednorodnym materiale w jednej odmianie, to tak niestety nie jest. Istnieje wiele różnic w parametrach pomiędzy wszystkimi typami FR-4, które są dostępne do wyboru dla projektantów i producentów.

Na szczęście dobrzy producenci radzą sobie z wieloma potencjalnymi problemami, które mogą pojawić się w procesie produkcyjnym i wybierają materiały, które to wytrzymają. Dzieje się tak dlatego, że są w stanie kontrolować i optymalizować proces tak, aby uniknąć komplikacji, a także są odpowiedzialni za zapewnienie działającej i niezawodnej płyty. Jeżeli którykolwiek z problemów związanych z rodzajem lub jakością laminatu dotyczy Twojego produktu, skonsultuj się z producentem w sprawie wyboru materiału spośród tych, które są kompatybilne z jego procesem.

Problemy nie kończą się na produkcji PCB, ponieważ podczas montażu i w docelowej aplikacji występują dodatkowe cykle cieplne. Zatem osoba montująca komponenty będzie musiała wziąć pod uwagę właściwości FR-4, na którym będzie montować komponenty. Warto omówić FR-4 używany przez producenta płytki PCB, aby upewnić się, że jest



Rysunek 5. Dodatkowe informacje o naszych płytach PCB

on kompatybilny i może wytrzymać cykle cieplne związane z montażem. Możesz też wybrać usługę, która wykona za Ciebie zarówno produkcję PCB, jak i montaż komponentów, aby kontrolować cały proces.

Po tym wszystkim weź pod uwagę temperaturę roboczą i lokalną na płycie. Czy te temperatury zbliżają się do T_g (każdy producent laminatu może podać margines poniżej T_g , poniżej którego należy pozostać)? Uwzględnij to wszystko przy wyborze laminatu i omów to ze swoimi dostawcami.

W Eurocircuits

Informacje o materiałach, z których korzystamy, znajdują się zawsze w zakładce *Pliki do pobrania*. Wygodnym miejscem, w którym możesz zobaczyć, jakich materiałów użyjemy do Twojej kompilacji, jest *Edytor kompilacji*. Widzimy tam poszczególne tuleje i prepregi, ich rodzaj i wysokość. Na **rysunku 5** pokazano kilka przykładów.

Eurocircuits

REKLAMA

Sięgnij po archiwalne wydania „ELEKTRONIKI PRAKTYCZNEJ”



Przesyłka
GRATIS

Zamów wygodnie na
www.UlubionyKiosk.pl



Produkcja wielowarstwowych płytek PCB krok po kroku

Nie ma czegoś takiego jak standardowa płytka PCB. Każda płytka drukowana ma unikalną funkcję zaprojektowaną dla konkretnej aplikacji. Produkcja PCB jest złożonym procesem składającym się z wielu etapów, a każdy z nich należy dostosować do wymagań projektu. W artykule opiszemy najważniejsze etapy produkcji wielowarstwowej płytki PCB wykonywanej w PCBWay.

Zamawiając płytki PCB w PCBWay otrzymujesz nie tylko profesjonalnie wykonany obwód drukowany, ale również jakość, która zwraca się z czasem. Gwarantuje to specyfikacja produktu i kontrola jakości, które są znacznie bardziej rygorystyczne niż w przypadku innych dostawców. Na diagramie z powyższego rysunku można zobaczyć, że proces produkcji w PCBWay jest wyjątkowy i nawet wykracza poza standard IPC.

1. Inżynieria przedprodukcyjna

Dane dostarczone przez klienta (najczęściej w postaci plików gerber) są używane do tworzenia danych produkcyjnych dla konkretnej płytki drukowanej (są to grafiki do procesów obrazowania i dane dotyczące wierceń do programów wierceń). Inżynierowie porównują wymagania z możliwościami, aby zapewnić zgodność, a także określić etapy procesu i powiązane kontrole. Żadne zmiany nie są dozwolone bez zgody Grupy PCBWay.

2. Cięcie cięcie laminatu platerowanego miedzią

Produkcja PCB rozpoczyna się od dużego kawałka arkusza materiału. Ze względu na ograniczenia sprzętu do produkcji PCB i możliwości



Fotografia 1.

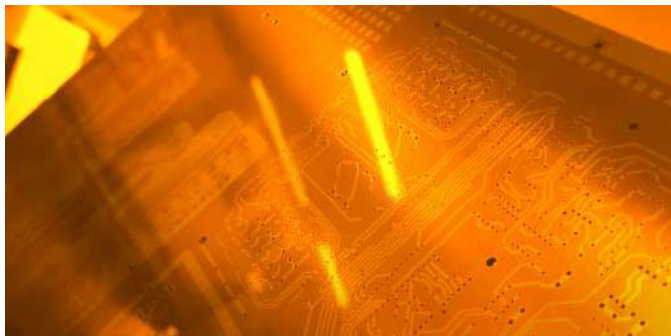
produkcyjnych, określone są wymagania dotyczące minimalnego i maksymalnego rozmiaru arkusza. Surowiec – laminat platerowany miedzią, należy przed rozpoczęciem produkcji przyciąć do rozmiaru przetworczego za pomocą automatycznej maszyny do cięcia – **fotografia 1**.

3. Nanoszenie schematu obwodu metodą naświetlania UV

Proces ten polega na przeniesieniu obrazu za pomocą specjalnego szablonu na powierzchnię płytki. Użycie światłoczułej folii i światła UV powoduje utwardzenie warstwy naświetlonej przez szablon – **fotografia 2**. Ten etap procesu wymaga idealnie czystych pomieszczeń.

4. Wytrawianie

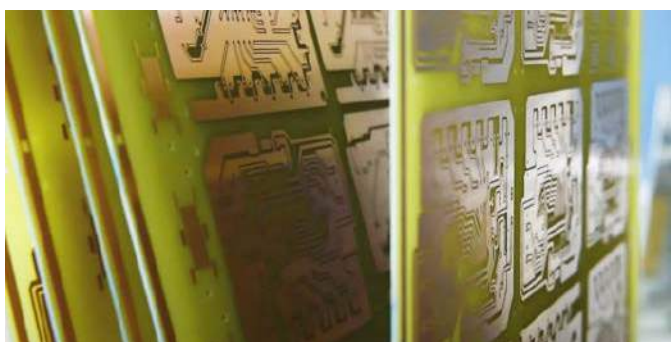
Chemiczne lub chemiczno-elektrolityczne usuwanie niepożądanych części materiału przewodzącego lub oporowego to wytrawianie. Pozostaje tylko miedziany obwód pasujący do projektu – **fotografia 3**.



Fotografia 2.

5. Automatyczna kontrola optyczna (AOI)

Kontrola obwodów na podstawie cyfrowych obrazów jest stosowana w celu sprawdzenia, czy obwody odpowiadają projektowi i czy są wolne od wad. Osiąga się to poprzez skanowanie arkuszy, a następnie przeszkoleni inspektorzy weryfikują wszelkie anomalie



Fotografia 3.



Fotografia 4.

wykryte w procesie skanowania – **fotografia 4**. Grupa PCBWay nie zezwala na naprawę otwartych obwodów.

6. Układanie i łączenie (laminowanie)

Poszczególne wytrawione warstwy projektu układa się razem z prepregiem zapewniającym izolację pomiędzy warstwami. Na górze i na dole stosu dodawana jest folia miedziana. Proces laminowania polega na umieszczeniu wszystkich warstw w ekstremalnej temperaturze 375 stopni Fahrenheit i ciśnieniu od 275 do 400 psi podczas laminowania za pomocą światłoczułej suchej warstwy ochronnej. PCB pozostawia się do utwardzenia w wysokiej temperaturze, ciśnienie jest powoli zwalniane, a następnie materiał jest powoli chłodzony.

REKLAMA

PCBWay

TWÓJ NAJLEPSZY PARTNER W PRODUKCJI PŁYTEK PCB!

DLACZEGO PCBWAY?

www.pcbway.com

service@pcbway.com



PŁYTKI PCB NAJWYŻSZEJ
JAKOŚCI



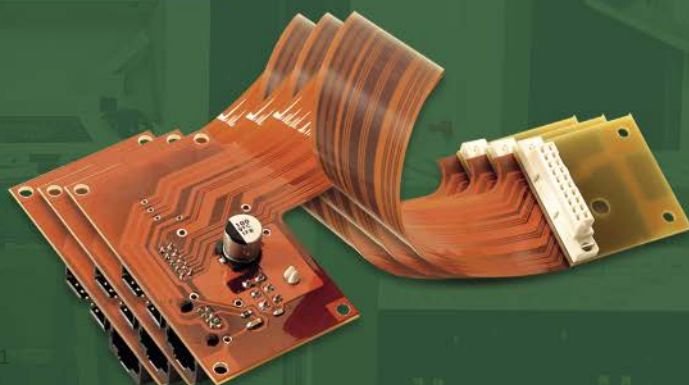
SZYBKI CZAS REALIZACJI

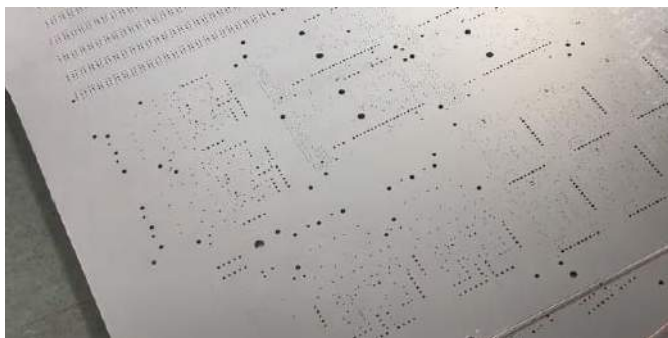


BEZPROBLEMOWE
SKŁADANIE ZAMÓWIENIA



CAŁODOBOWA
OBSŁUGA KLIENTA





Fotografia 5.

7. Wiercenie

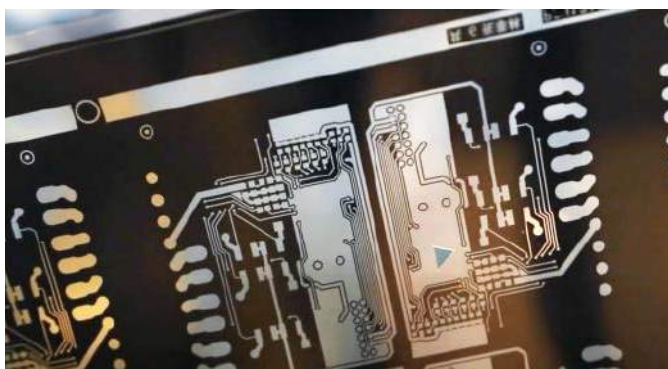
Musimy teraz wywiercić otwory, które następnie utworzą połączenia elektryczne w wielowarstwowej płycie PCB. Jest to proces wiercenia mechanicznego, który należy zoptymalizować, aby uzyskać pasowanie do wszystkich połączeń warstwy wewnętrznej – **fotografia 5**. W tym procesie panele można układać w stosy. Wiercenie można również wykonać za pomocą wiertarki laserowej.



Fotografia 6.

8. Osadzanie miedzi

Pierwszym krokiem w procesie galwanizacji jest chemiczne osadzanie bardzo cienkiej warstwy miedzi na ściankach otworów i całych panelach. Jest to złożony proces chemiczny, który musi być ściśle kontrolowany. Osadzanie miedzi nawet na niemetalowych ściankach otworów pozwoli uzyskać ciągłość elektryczną pomiędzy warstwami i przez otwory – **fotografia 6**. Aby zapewnić grubszą warstwę miedzi, zwykle od 5 do 8 μm stosuje się dodatkowe procesy galwanizacji/osadzania.



Fotografia 7.

9. Nanoszenie obwodów warstw zewnętrznych

Jest to etap podobny do wcześniejszego nanoszenia obwodu metodą naświetlania UV, ale z jedną główną różnicą – usuwamy suchą warstwę tam, gdzie chcemy zachować miedź – **fotografia 7**. Dzięki temu możemy nakładać dodatkowe powłoki miedzi na późniejszym etapie procesu. Ten proces przeprowadza się w pomieszczeniu o wysokiej czystości.



Fotografia 8.

10. Nanoszenie miedzi i cynowania

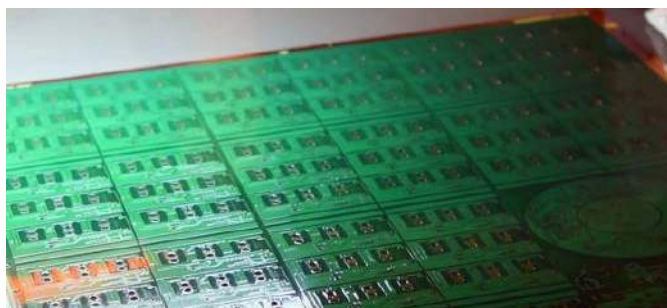
Drugi etap galwanizacji elektrolitycznej, w którym dodatkowe pokrycie osadza się w obszarach pozbawionych suchej warstwy (obwód elektryczny) – **fotografia 8**. Po pokryciu miedzią nakładana jest cyna, aby chronić platerowaną miedź.

11. Wytrawianie warstwy zewnętrznej

Zwykle jest to proces składający się z trzech etapów. Pierwszym krokiem jest usunięcie suchego filmu, drugi etap polega na wytrawieniu odsłoniętej/niepożądanego miedzi, podczas gdy osad cyny działa jako powłoka odporna na trawienie, chroniąc potrzebną nam miedź. Trzecim i ostatnim krokiem jest chemiczne usunięcie osadu cyny.

12. Kontrola AOI zewnętrznej warstwy

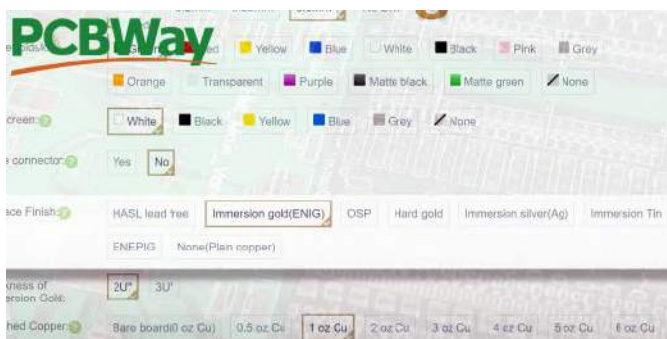
Podobnie jak w przypadku warstwy wewnętrznej, obrazowany i trawiony panel jest skanowany, aby upewnić się, że obwód odpowiada projektowi i jest wolny od defektów. Ponownie, zgodnie z wymaganiami PCBWay, naprawa otwartych obwodów nie jest dozwolona.



Fotografia 9.

13. Maska lutownicza

Preparat do maski lutowniczej nakładany jest na całą powierzchnię PCB – **fotografia 9**. Używając szablonów i światła UV, narażamy pewne obszary na działanie promieni UV, a obszary nienaświetlone są usuwane w procesie chemicznym – zazwyczaj są to obszary, które mają być użyte jako powierzchnie lutownicze. Pozostała maska lutownicza jest następnie całkowicie utwardzana i nanoszona jest warstwa opisowa. Ten etap procesu przeprowadza się w czystym pomieszczeniu.



Fotografia 10.



Fotografia 11.

14. Wykończenie powierzchni

Na odsłonięte obszary miedzi nakładane są różne wykończenia. Ma to na celu umożliwienie ochrony powierzchni i dobrą lutowalność. Różne wykończenia mogą obejmować pozłacanie, niklowanie, HASL, posrebrzanie itp. dostępne w różnych grubościach – **fotografia 10**. Zawsze przeprowadzane są testy grubości i lutowalności.

15. Cięcie

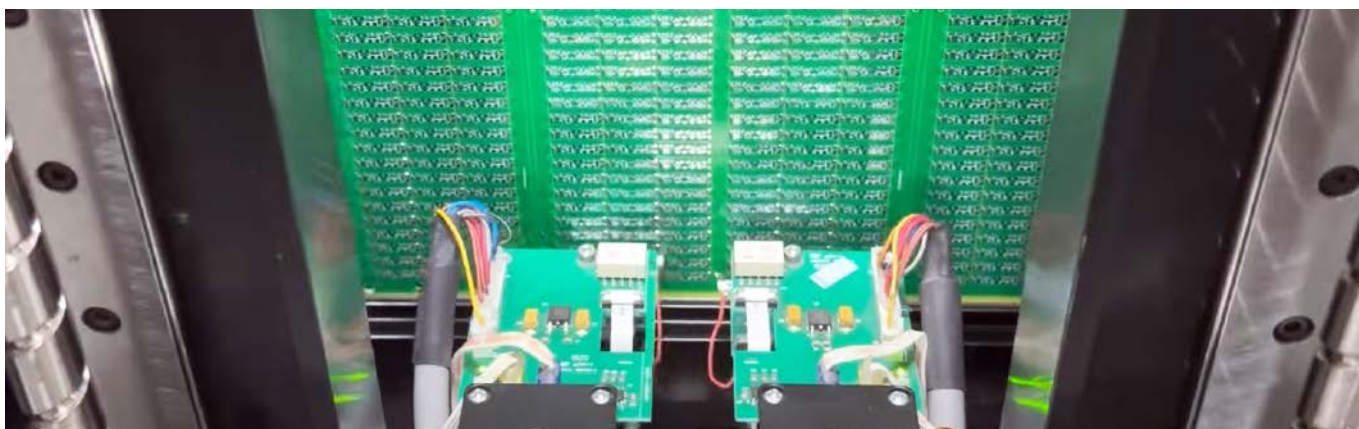
Proces cięcia paneli produkcyjnych nadaje płytkom określone rozmiary i kształty zgodne z projektem klienta zdefiniowanym w plikach gerber – **fotografia 11**. Dostępne są 3 główne opcje – punktacja, trasowanie lub wykrawanie. Wszystkie wymiary są mierzone na podstawie rysunku dostarczonego przez klienta, aby upewnić się, że panel ma prawidłowe wymiary.

16. Testy elektryczne

Ten etap służy do sprawdzania integralności ścieżek i połączeń otworów przelotowych. Pozwala sprawdzić, czy na gotowej płytce nie ma otwartych obwodów ani zwarc. Istnieją trzy metody testowania: latająca sonda do mniejszych projektów i 4-przewodowe testowanie Kelvina (do PCB dla przemysłu motoryzacyjnego lub lotniczego). Każdą płytkę drukowaną testujemy elektrycznie. Za pomocą testera latającej sondy sprawdzamy każdą sieć, aby upewnić się, że jest kompletna (nie ma przerw w obwodach) i nie ma zwarc z żadną inną siecią – **fotografia 12**.

17. Kontrola końcowa

Na ostatnim etapie produkcji zespół wnikliwych inspektorów poddaje każdą płytkę drukowaną ostatecznej, dokładnej kontroli. Następuje wizualne sprawdzenie płytki PCB pod kątem kryteriów akceptacji PCBWay.



Fotografia 12.



Fotografia 13.

Korzystanie z autoamtycznej inspekcji wizualnej – porównania PCB z plikiem gerber, ma większą prędkość sprawdzania niż ludzkie oko, ale nadal wymaga weryfikacji przez człowieka. Wszystkie zamówienia poddawane są także pełnej kontroli obejmującej wymiar, lutowalność itp.

18. Pakowanie

Płytki są pakowane przy użyciu materiałów zgodnych z wymaganiami PCBWay Packaging (ESD itp.), a następnie pakowane w pudełka przed wysyłką przy użyciu żądanego środka transportu – **fotografia 13**.

Podsumowanie

Na pierwszy rzut oka płytki PCB niewiele się różnią. Natomiast pod powierzchnią skupiamy się na wszystkich aspektach krytycznych dla trwałości i funkcjonalności płytek PCB. Klienci nie zawsze widzą różnicę, ale mogą być pewni, że PCBWay dokłada wszelkich starań, aby zapewnić swoim klientom płytki PCB spełniające najbardziej rygorystyczne standardy jakości.

Żadne zamówienie nie jest za małe ani za duże. Nasi klienci osiągają najlepszy możliwy czas wprowadzenia produktu na rynek i przewagę konkurencyjną, produkując płytki drukowane w sposób zrównoważony po najniższych kosztach całkowitych, dzięki naszym kompetencjom, dokładności dostaw i jakości produktu.

Pierwszą, standardową ofertą jest Quick-turn PCB, co oznacza, że możemy zaoferować małe ilości płytek PCB z możliwością szybkiej realizacji. Druga oferta – Advanced PCB, pokazuje to, co najlepsze, co PCBWay może zaoferować. Są to płytki PCB o pełnej specyfikacji, wysoce wyspecjalizowane, precyzyjne i gotowe do produkcji na dużą skalę. Technologia produkcji pozwala na realizację obwodów 64-warstwowych, z minimalnymi ścieżkami i przestrzeniami 2,5/2,5 milsa.

www.pcbway.com



Wielowarstwowe płytki drukowane,

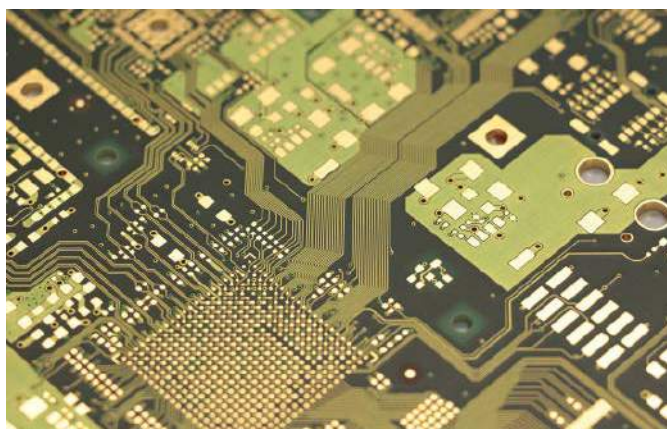
co każdy projektant wiedzieć powinien?

Świadome projektowanie obwodów wielowarstwowych wymaga od projektanta nie tylko doskonałego warsztatu – doświadczenia, umiejętności i sporej dozy inżynierskiej intuicji, ale także obszernej wiedzy z zakresu materiałoznawstwa elektronicznego, kompatybilności elektromagnetycznej oraz wymogów normalizacyjnych. W artykule opisujemy najważniejsze zagadnienia, z którymi każdy twórca płytek wielowarstwowych wcześniej czy później zetknie się w swojej praktyce.

W materiale pt. *Zagadnienia materiałowe w produkcji wielowarstwowych obwodów drukowanych* omówiliśmy szereg aspektów, związanych z parametrami rdzeni, prepregów i folii miedzianej, stosowanych do budowy stosu PCB. Nie mniej ważne, niż odpowiedni dobór materiałów, są również kwestie efektywnego użycia dostępnej liczby warstw czy też użycia różnych rodzajów przelotek. Te ostatnie są zresztą najczęściej dyskutowanym zagadnieniem podczas omawiania płyt typu HDI (**fotografia 1**) – to właśnie dzięki rozbudowanej technologii produkcji przelotek możliwe jest ciągle zwiększanie gęstości połączeń, nierzadko za pomocą najbardziej zaawansowanych technik, które do niedawna dostępne były tylko dla nielicznych światowych koncernów.

Budowa stosu z układowego punktu widzenia – warstwy sygnałowe, płaszczyzny masy i zasilania

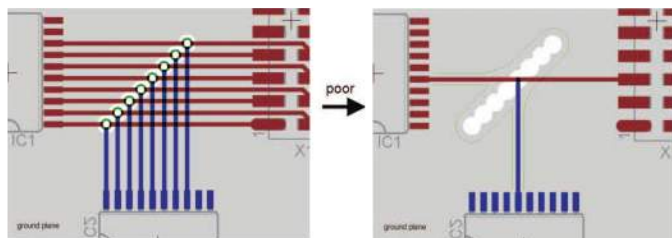
Określenie liczby warstw jest zwykle jedną z pierwszych czynności, które wykonuje każdy projektant płyty wielowarstwowej. Na tę decyzję wpływają nie tylko czynniki techniczne, ale także ekonomiczne – pomimo że zdecydowana większość współczesnych fabryk PCB



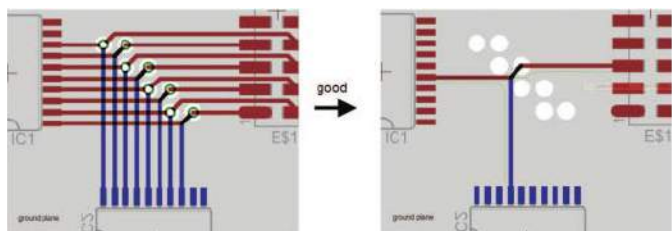
Fotografia 1. Przykład obwodu drukowanego klasy HDI
(https://t.ly/RLV_-)

oferuje możliwość wytworzenia przynajmniej podstawowych rodzajów obwodów 4-, 6- czy 8-warstwowych, to wciąż ich cena jest znacznie wyższa niż w przypadku obwodów 2-warstwowych. I jak zwykle bywa, także tutaj kompromis jest rozwiązaniem najlepszym – zbyt duża liczba warstw niepotrzebnie zwiększa koszty i wydłuża czas produkcji, zaś zbyt mała – ogranicza możliwości projektanta zarówno w zakresie podstawowego trasowania ścieżek, jak i optymalizacji projektu pod względem kompatybilności elektromagnetycznej, jakości sieci zasilających czy też osiągnięć w kwestii odprowadzania nadmiaru ciepła.

Niektóre zasady dobrej praktyki projektowej obowiązują niemal zawsze, niezależnie od stopnia złożoności urządzenia oraz liczby warstw. Elementem każdej płyty wielowarstwowej o znaczeniu krytycznym dla



Rysunek 1. Przykład błędnego ułożenia grupy przelotek sygnałowych; po lewej – layout płytki, po prawej – uwidoczniona nieciągłość płaszczyzny masy, wynikająca ze „złania się” antypadów w jedno, podłużne wycięcie (<https://t.ly/qRDsV>)

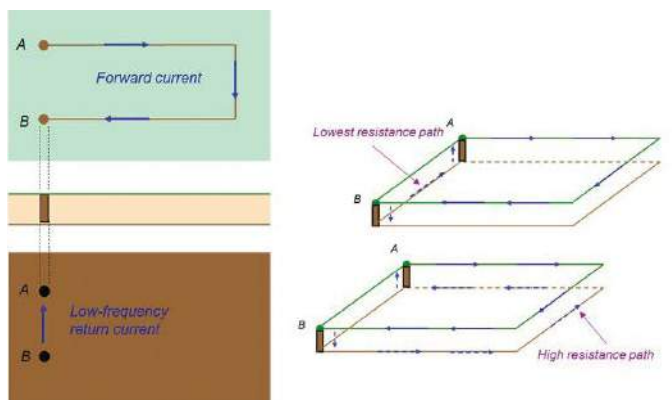


Rysunek 2. Poprawiony layout z rysunku 1 (<https://t.ly/qRDsV>)

poprawnego działania całości jest przynajmniej jedna, solidna płaszczyzna masy, pozbawiona (w miarę możliwości) jakichkolwiek większych wycięć. Zadań takiej warstwy jest wiele – pełni ona funkcję podstawowego ekranu oraz płaszczyzny odniesienia dla ścieżek paskowych i mikropaskowych, dystrybuuje potencjał „zerowy” dla wszystkich obwodów znajdujących się na płycie, a także przewodzi prądy powrotne, sprzężone z przebiegami przesyłanymi za pomocą ścieżek, znajdujących się na sąsiadujących warstwach sygnałowych. I właśnie ciągłość masy jest jednym z najistotniejszych zagadnień, poruszanych przez znawców tematyki EMC oraz integralności sygnałów, zwłaszcza w szybkich systemach cyfrowych.

Dobra masa to podstawa – znaczenie płaszczyzny masy dla EMC i integralności sygnałów

Przykładem typowego błędu, popełnianego niestety przez wielu projektantów, jest układanie przelotek sygnałowych w odstępach tak małych, że (przy przyjętych regułach DRC w zakresie odstępów izolacyjnych) antypady „zlewają się”, co skutkuje powstaniem podłużnego obszaru pozbawionego połączenia z masą (**rysunek 1**). Choć pozornie ów błąd nie powinien stanowić większego problemu, to w rzeczywistości taka nieciągłość sprawia, że prąd powrotny musi „omijać” cały rząd przelotek, co prowadzi do wytworzenia swego rodzaju „pętli”, obejmującej swoim zasięgiem niepotrzebnie spory obszar



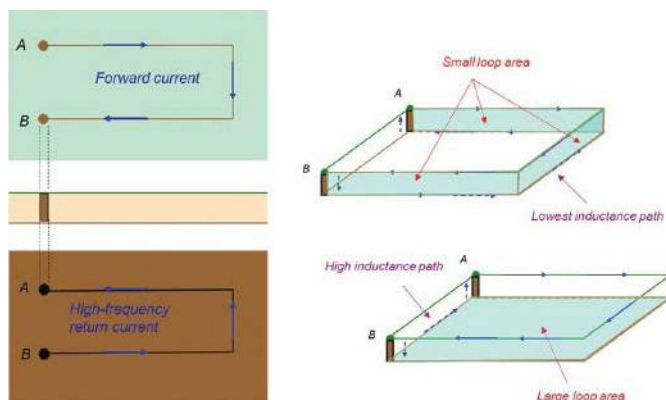
Rysunek 3. Prosty przykład obwodu ukazuje drogę o najniższej rezystancji, „preferowaną” przez prąd powrotny w zakresie od DC do około 100 kHz – sygnał jest przesyłany ścieżką na warstwie górnej z punktu A do B, zaś odpowiadający mu prąd powrotny płynie w warstwie dolnej, w prostej linii od punktu B do punktu A (<https://t.ly/5gMra>)

PCB. Tymczasem już niewielka zmiana, polegająca na nieznacznym rozsunięciu przelotek i ustawieniu ich w dwóch naprzemiennych rzędach, sprawia, iż prądy powrotne mogą swobodnie przepływać pomiędzy antypadami (**rysunek 2**), co wydatnie ogranicza długość ścieżki prądowej – a to wystarczy, by zmniejszyć emisję zakłóceń i zredukować poziom szumów, propagowanych za pośrednictwem masy.

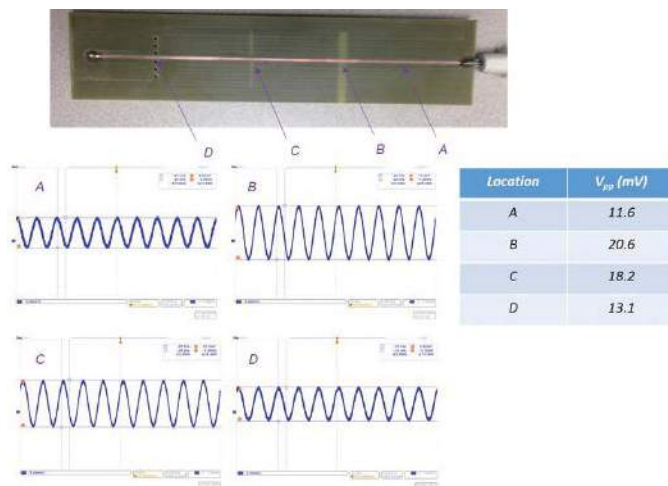
Opisane zjawisko warto zrozumieć dokładniej, gdyż dotyczy ono znacznie szerszego zakresu zagadnień projektowych niż tylko wspomnianych grup przelotek. Na **rysunku 3** pokazano prosty, teoretyczny model obwodu, złożonego z dwóch punktów A i B, znajdujących się pod napięciem stałym lub przemiennym o niskiej częstotliwości i połączonych za pomocą pętli przewodnika na górnej warstwie oraz masy na dolnej. Jak widać, choć trasa prądu na górnej warstwie jest wymuszona przez kształt ścieżki, to już na płaszczyźnie masy to ograniczenie nie istnieje – prąd powraca do punktu A drogą o najniższej rezystancji, czyli po prostu... wzdłuż odcinka B–A. W miarę wzrostu częstotliwości rośnie także znaczenie indukcyjności wypadkowej – preferowaną drogą prądu na płaszczyźnie masy będzie już trasa o najniższej indukcyjności, tj. taka, która – wraz z górną ścieżką sygnałową – zamyka pętlę o możliwie najmniejszym polu powierzchni. Takiej właśnie sytuacji odpowiada **rysunek 4** – z uwagi na znikomą grubość dielektryka pole będzie najmniejsze, gdy droga powrotna będzie znajdowała się tuż pod ścieżką górną.

Słuszność powyższego opisu można zresztą bardzo łatwo udowodnić prostym eksperymentem. Na **rysunku 5** pokazano oscylogramy, zarejestrowane w czterech punktach prostej płytki testowej, która w płaszczyźnie masy ma trzy przeszkody (widoczne na fotografii w punktach B, C i D). Gęstość prądu zmierzono w sposób pośredni, poprzez przyłożenie sondy pola bliskiego (typu H, czyli rejestrującej pole magnetyczne), podłączonej do wejścia oscyloskopu. Najsilniejszy sygnał zarejestrowano w punkcie B – tam prąd musi „przeciskać się” przez wąski pasek miedzi, pozostawiony tylko na jednym brzegu płytki. Przeszkoda umieszczona w punkcie C pozwala na przepływ prądu po obydwu stronach ubytku miedzi (tutaj amplituda sygnału jest już nieco niższa), zaś rząd otworów w punkcie D odpowiada opisanej wcześniej grupie przelotek (w tym przypadku jednak pozbawionych antypadów, a więc nieblokujących całkowicie przepływu prądu pomiędzy otworami – z tego względu amplituda sygnału jest tylko nieznacznie wyższa niż w punkcie referencyjnym A).

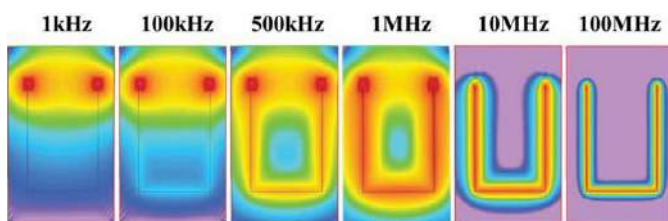
Warto jeszcze rzucić okiem na wyniki symulacji rozkładu gęstości prądu, pokazane na **rysunku 6** – im wyższa częstotliwość, tym silniej prąd „skupia się” wokół ścieżki sygnałowej. Jak widać, już dla 100 kHz część prądu zaczyna „uciekać” z linii najmniejszego oporu i tworzyć „cień” ścieżki sygnałowej, zaś powyżej 1 MHz ów „cień” dominuje w coraz wyższym stopniu. Pokazane wyniki jasno pokazują, jak wielkie znaczenie dla obwodów wielowarstwowych ma solidna,



Rysunek 4. Sytuacja analogiczna jak na rysunku 2 z tą różnicą, że dotyczy prądu przemiennego o wysokiej częstotliwości. W tym przypadku preferowana jest ścieżka o najmniejszej indukcyjności – pole pętli prądowej jest najmniejsze, jeżeli droga prądu powrotnego przebiega bezpośrednio pod ścieżką sygnałową (<https://t.ly/5gMra>)



Rysunek 5. Wyniki eksperymentu, w którym sygnał z generatora funkcyjnego doprowadzono do obciążenia w postaci pętli, złożonej z prostej ścieżki na warstwie górnej, zwartej z płaszczyzną masy na warstwie dolnej. Pośredniego pomiaru gęstości prądu dokonano sondą pola bliskiego typu H w czterech punktach: przy wejściu sygnałowym oraz w trzech różnych punktach nieciągłości masy: B – szerokie wycięcie, dochodzące z jednej strony do brzegu PCB; C – wąskie wycięcie położone centralnie na szerokości płytki; D – szereg niewielkich otworów, pomiędzy którymi znajdowały się wąskie przejścia miedzi. W tabeli podano amplitudę sygnałów, zarejestrowanych za pomocą oscyloskopu (<https://t.ly/5gMra>)



Rysunek 6. Symulacja rozkładu gęstości prądu powrotnego pomiędzy dwoma punktami obwodu dla sześciu różnych częstotliwości (<https://t.ly/5gMra>)

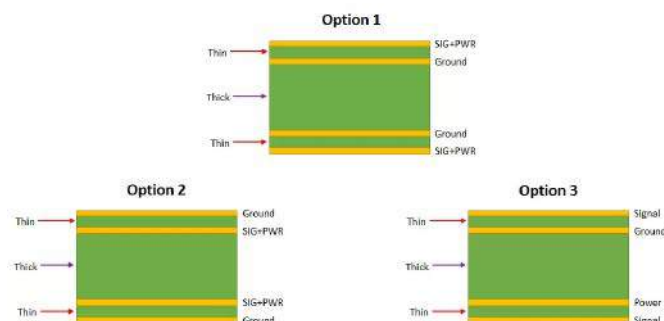
jednolita płaszczyzna masy – swobodny przepływ prądów powrotnych pod ścieżkami prowadzącymi szybkie sygnały (zwłaszcza cyfrowe) pozwala nie tylko znacząco zmniejszyć natężenie zakłóceń RFI (im mniejsza pętla prądowa, tym słabsza emisja i mniejszy wpływ na pobliskie obwody), ale także może poprawić integralność sygnałów dzięki zachowaniu najmniejszej możliwej indukcyjności przy zadanej geometrii i długości ścieżki sygnałowej.

Konfiguracja stosu warstw w zależności od aplikacji

Z matematycznego punktu widzenia wraz ze zwiększaniem liczby warstw PCB drastycznie rośnie także liczba możliwych do zastosowania konfiguracji ich ułożenia. W praktyce zdecydowana większość płyt wielowarstwowych jest oparta na pewnych powtarzalnych schematach, sensownych z punktu widzenia czterech głównych czynników:

- integralności sygnałów,
- efektywności sieci zasilających,
- kontroli impedancji,
- kompatybilności elektromagnetycznej.

W przypadku PCB 4-warstwowych warto rozważenia są trzy główne koncepcje, pokazane schematycznie na **rysunku 7**. Opcja pierwsza używa dwóch wewnętrznych warstw do utworzenia płaszczyzny masy, zaś wszystkie ścieżki sygnałowe oraz zasilające są rozdysponowane na warstwach zewnętrznych. Taka konstrukcja pozwala na prowadzenie szybkich linii transmisyjnych (np. USB, HDMI czy LVDS), obwodów RF czy też wszelkiego rodzaju innych interfejsów, wymagających wysokiej częstotliwości pracy i dokładnej kontrolowanej

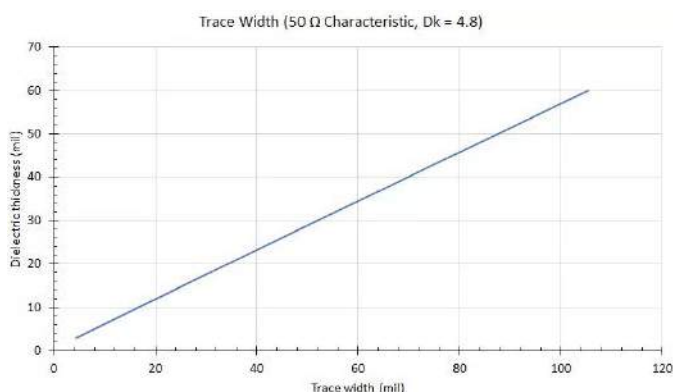


Rysunek 7. Trzy najczęściej wykorzystywane konstrukcje stosu 4-warstwowego (<https://t.ly/Pvj16>)

impedancji. Ponieważ w wielu projektach stosowana jest budowa stosu typu foil build (więcej informacji na ten temat można znaleźć w artykule pt. *Zagadnienia materiałowe w produkcji wielowarstwowych obwodów drukowanych*), odległość pomiędzy płaszczyzną masy a warstwą zewnętrzną po tej samej stronie płytki drukowanej jest zwykle niewielka (i określona tylko przez finalną grubość prepregu). To zaś umożliwia wydatną redukcję szerokości ścieżek, która – dla ustalonej impedancji charakterystycznej i stałej dielektrycznej prepregu – zmienia się wprost proporcjonalnie do grubości dielektryka (patrz **rysunek 8**). Niestety, stos wg opcji 1 ma także swoje wady – wymusza konieczność zmieszczenia wszystkich sygnałów i ścieżek zasilania na dwóch warstwach, nie ekranuje linii mikropaskowych przed zewnętrznymi zakłóceniami RFI oraz nie zapewnia tłumienia zaburzeń elektromagnetycznych, promieniowanych przez obwody znajdujące się na zewnętrznych warstwach PCB. Doskonała jest natomiast separacja pomiędzy ścieżkami, znajdującymi się na warstwach *top* i *bottom*.

Opcja 2 według rysunku 7 jest dokładnym odwróceniem poprzednio opisanej koncepcji – tym razem sygnały i zasilanie znajdują się wewnątrz „kanapki”, utworzonej przez zewnętrzne płaszczyzny masy. Takie rozwiązanie zapewnia perfekcyjne osiągi PCB pod względem emisji i podatności EMI, ale naraża linie sygnałowe na przesłuchy z warstwy leżącej naprzeciwko. Co ważne, także w tym przypadku istnieje możliwość prowadzenia szybkich ścieżek sygnałowych o kontrolowanej impedancji na dwóch warstwach, ale – niestety – z pewną dość istotną wadą: każde wyprowadzenie elementu SMD (oprócz pinów masy) musi być wprowadzone w głąb PCB za pomocą przelotki, co dodaje pewną indukcyjność i samo w sobie pogarsza nieco integralność sygnałów.

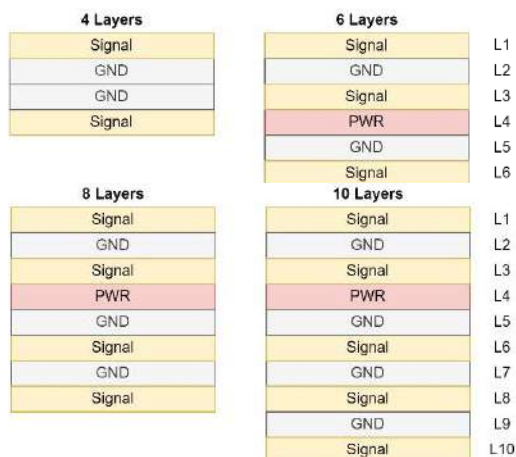
Opcja 3 jest natomiast najlepszym wyborem w przypadku urządzeń, w których liczba połączeń o kontrolowanej impedancji jest niewielka – wszelkie sygnały niskoczęstotliwościowe lub DC można zatem prowadzić na warstwie zewnętrznej przylegającej do płaszczyzny zasilania, a warstwę leżącą po stronie masy warto przeznaczyć



Rysunek 8. Zależność grubości dielektryka od szerokości ścieżki, przy ustalonej docelowej impedancji charakterystycznej równej 50 Ω i stałej dielektrycznej równej 4,8 (https://t.ly/tM_GI)



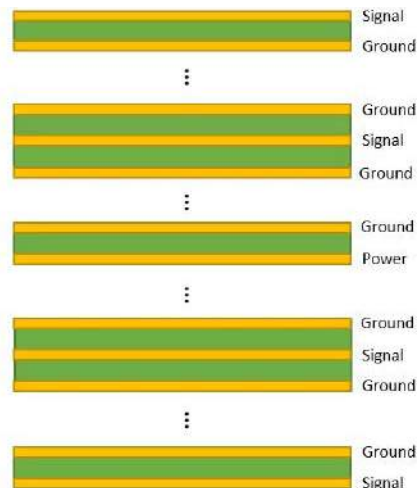
Rysunek 9. Przykład płaszczyzny zasilania, podzielonej na szeregi sekcji o różnym przeznaczeniu (np. zasilanie części analogowej, cyfrowej i RF) oraz różnych poziomach napięcia (<https://t.ly/DnLRY>)



Rysunek 10. Przykładowe konfiguracje stosów płyt 4-, 6-, 8- i 10-warstwowych (<https://t.ly/Padd5>)

do wszystkich bardziej wymagających połączeń. Ta konfiguracja stosu lepiej sprawdza się także w przypadku układów o większym poborze mocy, gdyż duża powierzchnia płaszczyzny zasilania pozwala na znaczne podwyższenie amperażu w porównaniu do (nawet szerokich) ścieżek, które prowadzone byłyby w warstwie sygnałowej. Mało tego – płaszczyznę zasilania można podzielić na kilka mniejszych obszarów, odpowiedzialnych za dystrybucję różnych napięć zasilania do poszczególnych części obwodu (**rysunek 9**).

W przypadku liczby warstw przekraczającej 4 zasady przyporządkowania poszczególnych płaszczyzn do określonych celów są analogiczne, do pokazanych do tej pory. I tak każda warstwa sygnałowa powinna być w miarę możliwości oddzielona od pozostałych płaszczyzną masy, zaś płaszczyzna zasilania oraz jedna z mas powinny tworzyć parę, przedzieloną tylko cienką warstwą dielektryka (prepregu lub rdzenia o niewielkiej grubości). To ostatnie jest szczególnie ważne, zważywszy na doskonałe właściwości pojemności rozproszonej, tworzącej swego rodzaju jeden wielki kondensator odprowadzający o powierzchni całej płytki. Takie rozwiązanie doskonale poprawia właściwości obwodu wielowarstwowego w zakresie integralności sygnałów oraz emisji RFI. Porównanie klasycznych stosów 4-, 6-, 8- i 10-warstwowego pokazano na **rysunku 10**, natomiast ogólną zasadę ekstrapolacji opisanych powyżej reguł ilustruje **rysunek 11**.



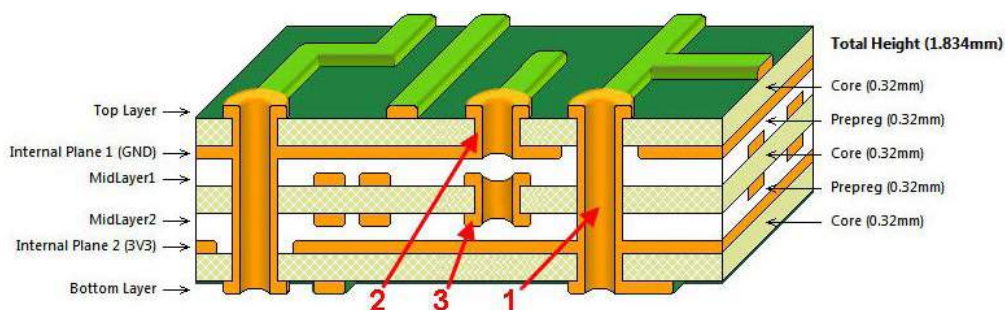
Rysunek 11. Strategia grupowania warstw w przypadku płyt o 12+ warstwach (<https://t.ly/USyCe>)

Rodzaje przelotek

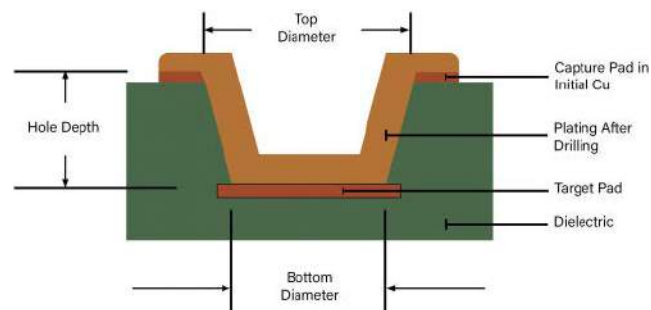
Przelotki stanowią jedno z fundamentalnych zagadnień w projektowaniu płyt wielowarstwowych. To właśnie szeroki wachlarz możliwości technologicznych, oferowanych przez najbardziej zaawansowanych producentów PCB w zakresie wytwarzania różnych typów przelotek, umożliwia ciągłe przesuwanie granic miniaturyzacji i sukcesywne zwiększanie gęstości połączeń.

Na **rysunku 12** pokazano schematycznie trzy podstawowe rodzaje przelotek, stosowanych w standardowych płytach wielowarstwowych. Przelotki określane mianem PTH to zwyczajne otwory metalizowane, przechodzące przez cały stos PCB – wykonywane dokładnie tak, jak otwory w padach elementów THT. Przelotki ślepe zaczynają się na jednej z warstw zewnętrznych i kończą wewnątrz stosu, zaś przelotki zagrzebane są całkowicie ukryte wewnątrz PCB. Zaletą tych dwóch ostatnich jest możliwość zaoszczędzenia miejsca na jednej, a nawet dwóch warstwach zewnętrznych, co niebawem ułatwia trasowanie ścieżek i rozmieszczanie komponentów na obszarach płytki o zwiększonej gęstości upakowania podzespołów.

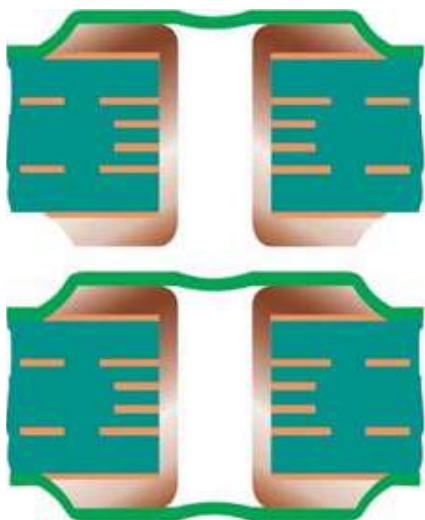
Kolejnym rodzajem przelotek, bez którego trudno wyobrazić sobie współczesną elektronikę, są tzw. mikroprzelotki – norma IPC-T-50M definiuje je jako ślepe przelotki o maksymalnym stosunku



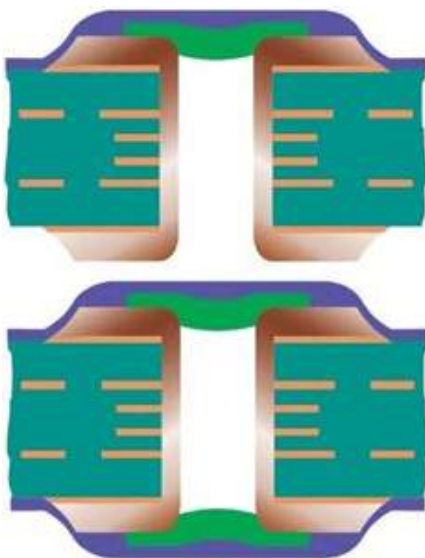
Rysunek 12. Trzy podstawowe rodzaje przelotek: 1 – przelotka typu PTH (Plated Through Hole), przechodząca przez cały stos PCB; 2 – przelotka ślepa; 3 – przelotka zagrzebana (<https://t.ly/FmAH4>)



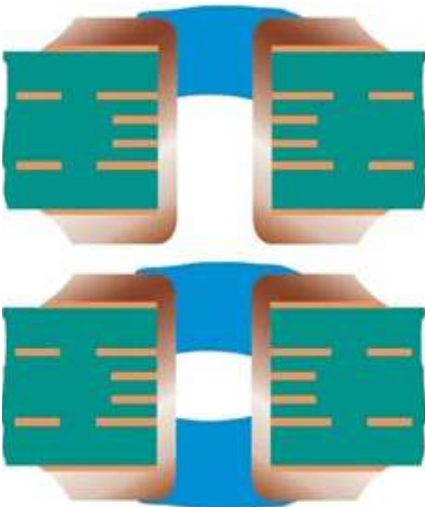
Rysunek 13. Definicja wymiarów mikroprzelotki (<https://t.ly/IF9Ww>)



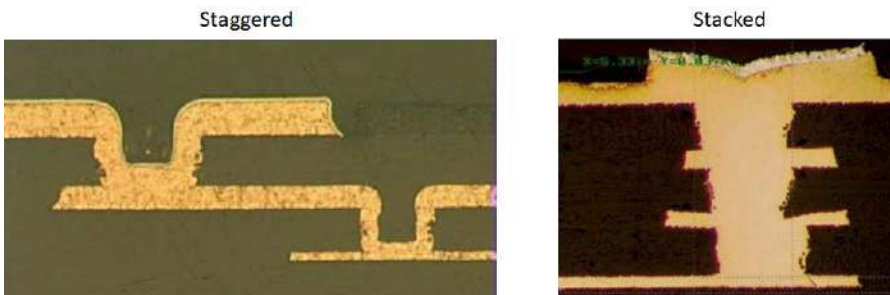
Rysunek 14. Przelotki typu I wg IPC-4761 (na górze – typ I-A, na dole – typ I-B) – <https://t.ly/YkoVX>



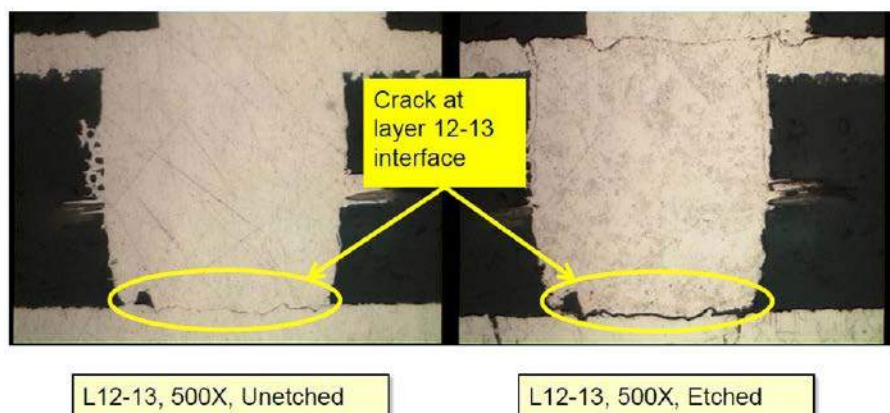
Rysunek 15. Przelotki typu II wg IPC-4761 (na górze – typ II-A, na dole – typ II-B) – <https://t.ly/YkoVX>



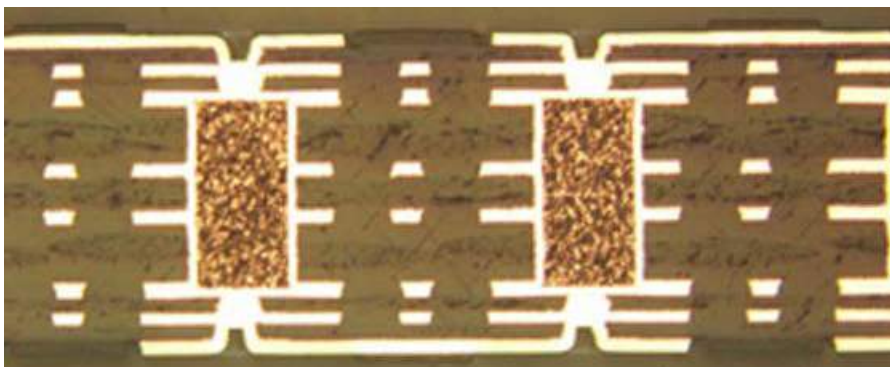
Rysunek 16. Przelotki typu III wg IPC-4761 (na górze – typ III-A, na dole – typ III-B) – <https://t.ly/YkoVX>



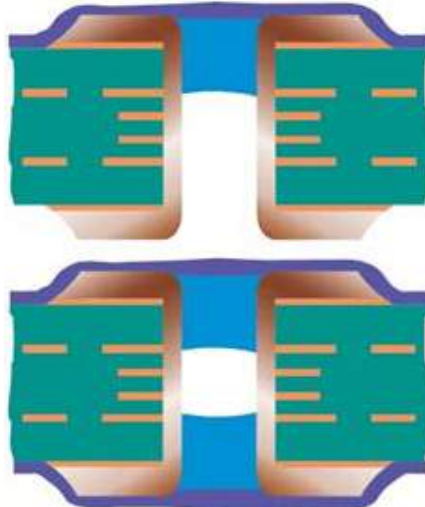
Fotografia 2. Przekroje mikroprielotek przesuniętych (staggered) i piętrowych (stacked) – <https://t.ly/5ztQ2>



Fotografia 3. Przykładowe uszkodzenie mikroprielotki – pęknięcie pomiędzy dolną powierzchnią metalizacji a padem docelowym (<https://t.ly/YhYmS>)



Fotografia 4. Przekrój mikroprielotek piętrowych, osadzonych na wypełnionych przelotkach zagrzebanych (https://t.ly/r5Z_)



Rysunek 17. Przelotki typu IV wg IPC-4761 (na górze – typ IV-A, na dole – typ IV-B) – <https://t.ly/YkoVX>

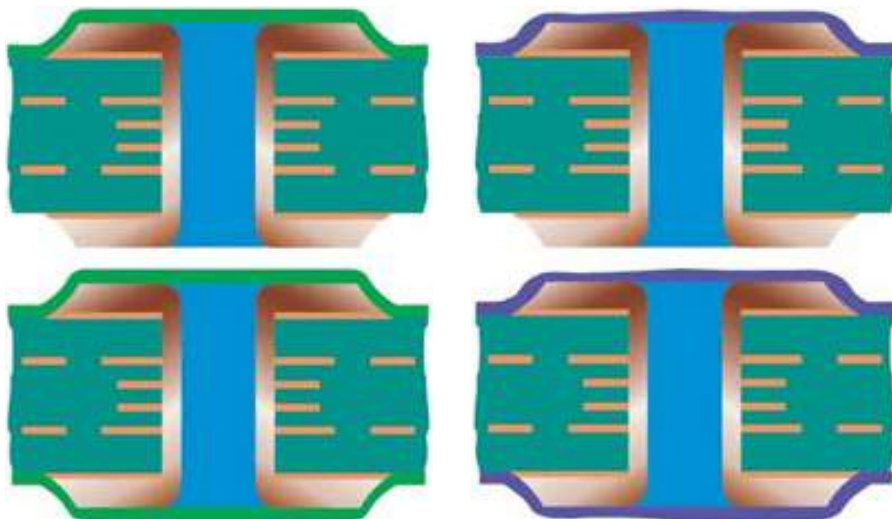
wymiarów (tzw. *aspect ratio*) równym 1:1 (tj. głębokość: średnica zewnętrzna – patrz rysunek 13), kończące się na padzie docelowym i mające głębokość nieprzekraczającą 0,25 mm (licząc od folii miedzianej, od której przelotka się zaczyna, aż do powierzchni padu docelowego). Podobną definicję znajdujemy także w standardzie IPC-6012. Warto zwrócić uwagę, że mikroprielotki – w odróżnieniu od klasycznych przelotek PTH, ślepych oraz zagrzebanych – nie są łączone z dolną warstwą za pomocą wykonanego w niej otworu, ale niejako „przyklejają się” do pełnego padu swoim dnem. Taka konstrukcja ma swoje zalety – pozwala wykonywać naprawdę miniaturowe połączenia pomiędzy sąsiadującymi ze sobą warstwami, zaś połączenie kolejnych warstw jest możliwe z użyciem trzech technik: przelotek typu *skip* (tj. kończących się nie na sąsiadującej warstwie, lecz na kolejnej,

z pominięciem warstwy pośredniczącej), *staggered* (kolejna przelotka zaczyna się na padzie, połączonym z padem docelowym przelotki leżącej o jedną warstwę wyżej) oraz *stacked* (dwie lub więcej przelotek jest ustawionych jedna na drugiej w postaci konstrukcji piętrowej) – patrz **fotografia 2**. Niestety, pomimo niewątpliwych zalet, mikroprzelotki mają także swoje wady – zauważono bowiem, że ten typ łączenia warstw jest podatny na uszkodzenia termomechaniczne, powodujące oderwanie metalizacji przelotki od padu docelowego, co można zobaczyć na **fotografii 3**. Oprócz zachowania prawidłowych warunków obróbki płyt bazujących na mikroprzelotkach, na niezawodność wpływa także liczba przelotek ustawionych piętrowo – zalecane są stosy co najwyżej dwóch przelotek, zaś w przypadku konieczności połączenia większej liczby warstw warto posilkować się konstrukcją mieszaną (*stacked-staggered*).

Warto także dodać, że mikroprzelotki mogą współpracować również ze standardowymi przelotkami zagrzebanymi, przy czym te ostatnie zawsze muszą zostać odpowiednio wypełnione – przykład takiej konstrukcji można zobaczyć na **fotografii 4**.

Kwestia wypełniania lub zabezpieczania przelotek z zewnątrz jest resztą znacznie szersza i obejmuje siedem klas głównych, określonych w normie IPC-4761.

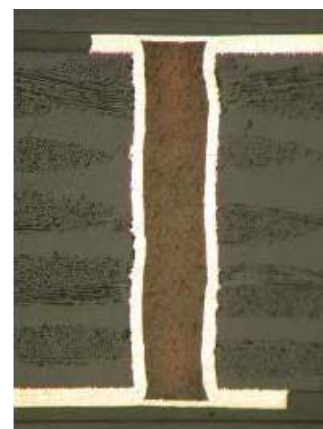
- **Przelotki typu I wg IPC-4761 (rysunek 14)** są wykonywane z użyciem tzw. tentingu, czyli pokrywania soldermaską z jednej lub obu stron; należy pamiętać, że nie jest to całkowite zabezpieczenie, a jedynie zgrubna osłona brzegów (pierścienia przelotki) – całkowite zatkanie otworu jest możliwe tylko dla przelotek bardzo małych (zwykle poniżej 0,3 mm), zaś większe otwory powodują wpływanie (stosunkowo rzadkiej) soldermaski do wnętrza i częściowe pokrywanie metalizacji. Jednostronne pokrycie (typ I-A) jest niezalecane z uwagi na ryzyko „pułapkowania” zanieczyszczeń chemicznych (np. resztek topnika o bardziej agresywnych właściwościach), mogących prowadzić do korozji metalizacji. Lepszą metodą jest zatem dwustronny tenting (typ I-B).
- **Przelotki typu II wg IPC-4761 (rysunek 15)** – w tym przypadku przelotki są pokrywane z jednej (II-A) lub dwóch stron (II-B) nie tylko podstawową soldermaską (tzw. suchą), ale także dodatkową warstwą soldermaski „mokrej” – takie rozwiązanie zapewnia lepszą osłonę wnętrza otworu przed czynnikami zewnętrznymi (oczywiście tylko w przypadku II-B, gdyż II-A nie jest rekomendowany z opisanego wcześniej powodu).
- **Przelotki typu III wg IPC-4761 (rysunek 16)** – tzw. przelotki zatykane (III-A, III-B) za pomocą dielektryka, np. epoksydu. Takie zabezpieczenie (zastosowane dwustronnie) całkowicie redukuje ryzyko penetracji przelotki przez lutowie lub zanieczyszczenia.
- **Przelotki typu IV wg IPC-4761 (rysunek 17)** – modyfikacja przelotek typu III z użyciem dodatkowego pokrycia za pomocą soldermaski mokrej (IV-A, IV-B).



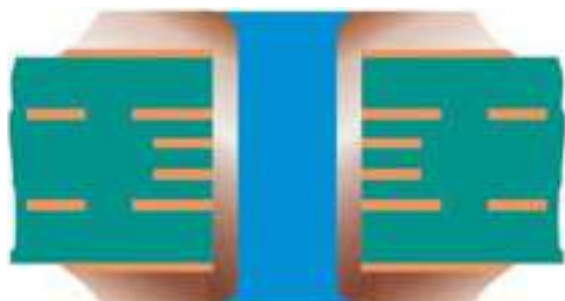
Rysunek 19. Przelotki typu VI wg IPC-4761 (na górze – typ VI-A, na dole – typ VI-B)
<https://t.ly/YkoVX>

- **Przelotki typu V wg IPC-4761 (rysunek 18)** – w tym przypadku światło otworu jest wypełniane dielektrykiem na całej długości.
- **Przelotki typu VI wg IPC-4761 (rysunek 19)** – modyfikacja przelotek typu V z dodatkowym pokryciem soldermaską suchą lub mokrą.
- **Przelotki typu VII wg IPC-4761 (rysunek 20)** – najdroższa w realizacji i najbardziej złożona technologicznie metoda, polegająca na wypełnieniu wnętrza przelotki, a następnie zastosowaniu metalizacji na całej jej średnicy.

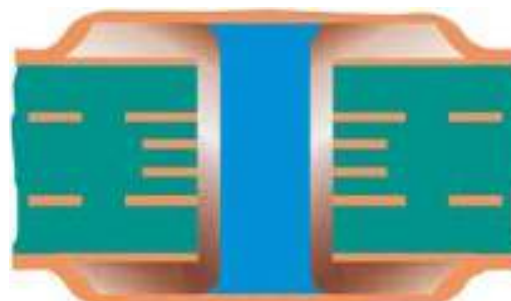
Warto dodać, że przelotki mogą być wypełniane nie tylko dielektrykiem, ale także przewodzącym polimerem (zwykle na bazie żywicy epoksydowej), a nawet... miedzią. Ten ostatni rodzaj wypełnienia jest szczególnie cenny z uwagi nie tylko na znaczną poprawę warunków elektrycznych (obniżenie rezystancji, wzrost obciążalności prądowej), ale także doskonałe parametry termiczne. Dlatego też przelotki wypełnione miedzią stanowią idealne rozwiązanie w przypadku chłodzenia elementów dużej mocy, np. układów scalonych SMD z padem termicznym. Należy pamiętać, że zastosowanie przelotek w padzie (ang. *via-in-pad*) jest możliwe tylko wtedy, gdy zostaną one uprzednio poddane tzw. planaryzacji, czyli procesowi polegającemu na mechanicznej obróbce powierzchni zewnętrznej przelotki tak, by licowała ona z powierzchnią pozostałych padów. Przekrój tak wykonanej przelotki można zobaczyć na **fotografii 5**.



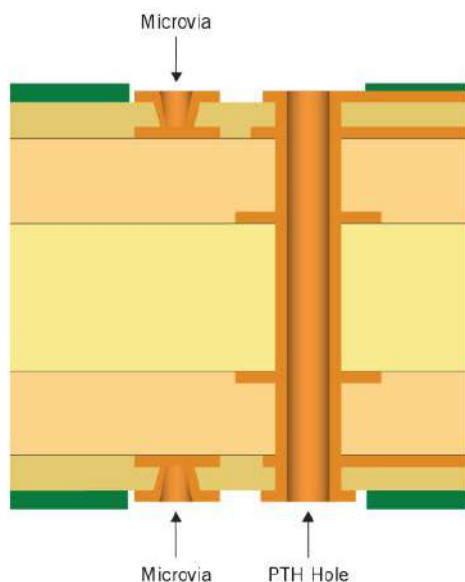
Fotografia 5. Przekrój zatkanej przelotki typu VII z planaryzacją
https://t.ly/_ykcd



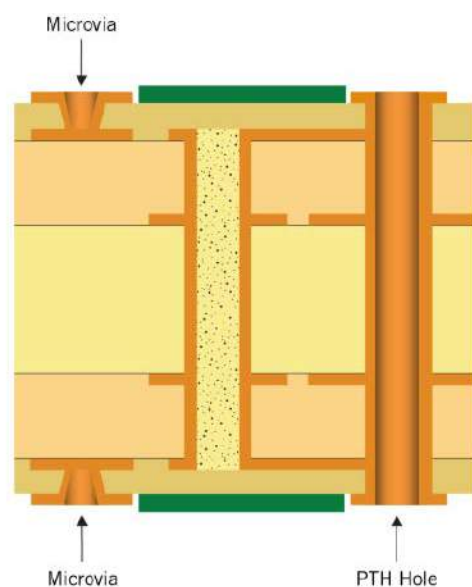
Rysunek 18. Przelotka typu V wg IPC-4761 – <https://t.ly/YkoVX>



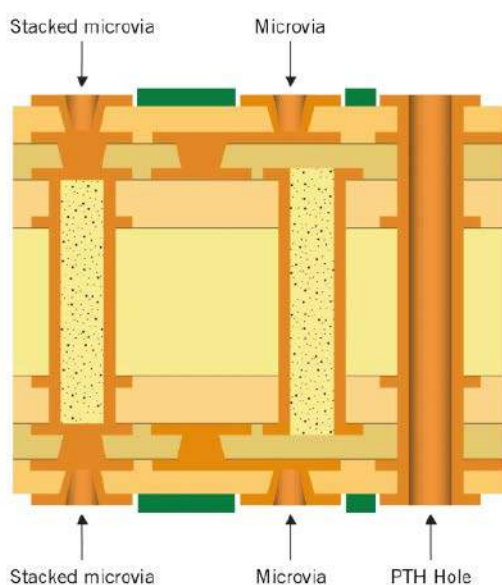
Rysunek 20. Przelotka typu VII wg IPC-4761 – <https://t.ly/YkoVX>



Rysunek 21. Struktura płytki HDI typu I wg IPC-2226 (https://t.ly/3_0oR)



Rysunek 22. Struktura płytki HDI typu II wg IPC-2226 (https://t.ly/3_0oR)



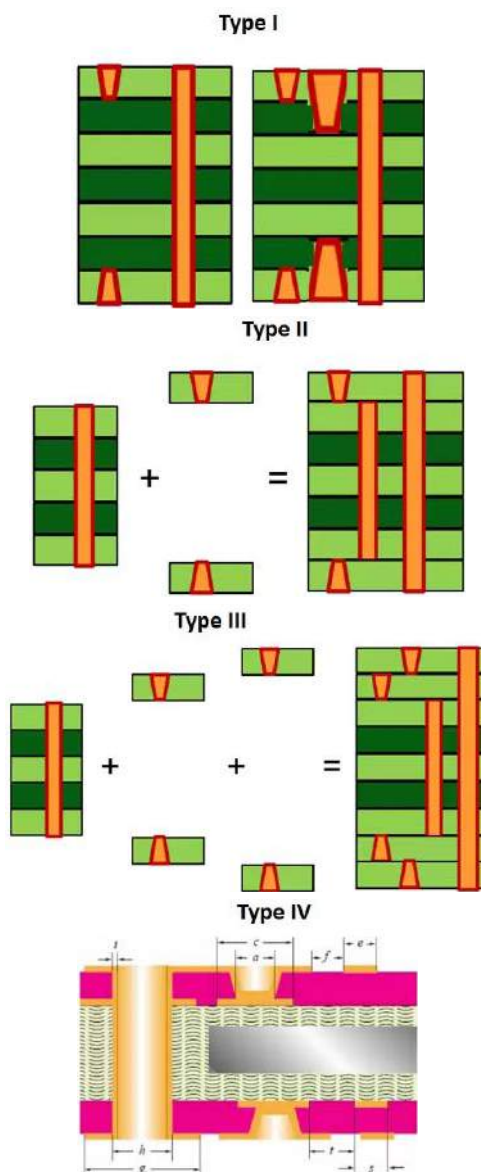
Rysunek 23. Struktura płytki HDI typu III wg IPC-2226 (https://t.ly/3_0oR)

Obwody HDI

Podstawowa definicja obwodów HDI (ang. *High Density Interconnect*) określa je jako płytki drukowane o gęstości połączeń (przypadającej na jednostkę powierzchni) wyższej w porównaniu do konwencjonalnych PCB. W rzeczywistości HDI to synonim nowoczesnych, wysoce zaawansowanych obwodów, wytwarzanych z użyciem technik dalece wykraczających poza ogólne standardy produkcji PCB. Typowym widokiem, spotykanym na płytkach HDI, są ultraminiaturowe komponenty pasywne w rozmiarach 0201, 01005, a nawet mniejszych (**fotografia 6**) oraz układy scalone w obudowach BGA, CSP itp.

Norma IPC-2226 definiuje sześć odmian obwodów HDI, zaś główną rolę w rozróżnieniu poszczególnych typów PCB grają przełotki. Trzy podstawowe typy obwodów (I, II, III) mają następującą konstrukcję:

- **struktura HDI typu I wg IPC-2226 (rysunek 21)** dopuszcza stosowanie przełotek typu PTH oraz mikroprzełotek na warstwach zewnętrznych,
- **struktura HDI typu II wg IPC-2226 (rysunek 22)** stanowi rozszerzenie typu I o przełotki zagrzebane, wykonywane metodą standardową,
- **struktura HDI typu III wg IPC-2226 (rysunek 23)** rozszerza typ II o mikroprzełotki zagrzebane, zaczynające się na warstwach leżących bezpośrednio pod warstwami zewnętrznymi (a zatem wchodzi tutaj w grę także przełotki typu *stacked* i *staggered*).



Rysunek 24. Schematyczne porównanie procesów konstrukcji płyt HDI typu I...III oraz budowa PCB typu IV wg IPC-2226 (<https://t.ly/00QZk>)



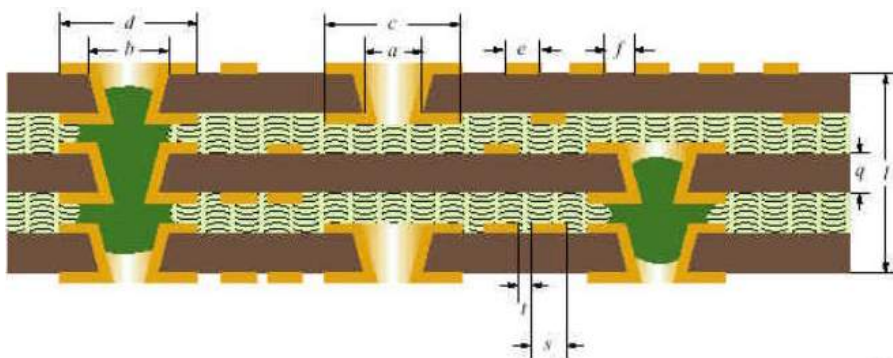
Fotografia 6. Porównanie wielkości komponentów SMD w rozmiarach 0402, 0201, 01005 oraz 008004. Ten ostatni to najmniejszy na świecie dławik ferrytowy, opracowany przez firmę Murata – jego wymiary to zaledwie 0,25×0,125 mm (<https://t.ly/7jRNd>)

Choć taki podział może wydawać się nieco sztuczny, to w istocie ma on naprawdę jasny sens praktyczny. Aby lepiej zrozumieć filozofię trzech opisanych rodzajów PCB, warto przyrzeć się schematom procesów produkcji, pokazanym na **rysunku 24**. Płytki typu I są laminowane w jednym, wspólnym procesie – otwory mikroprzelotek mogą być wykonane w warstwach zewnętrznych jeszcze przed sprasowaniem całego stosu, zaś otwory PTH – już po laminacji. W przypadku typu II stosowany jest już proces laminacji sekwencyjnej – najpierw powstaje zestaw warstw o numerach od 2 do (n-1), potem mogą nastąpić wiercenia przelotowe (które staną się podstawą przelotek zagrzebanych), a na samym końcu dodawane są warstwy zewnętrzne z otworami pod mikroprzelotki. Typ III korzysta z tej samej filozofii, ale laminacja odbywa się w trzech sekwencjach, stąd możliwość wytworzenia także mikroprzelotek zagrzebanych, rozpoczynających się na warstwach 2 i (n-1).

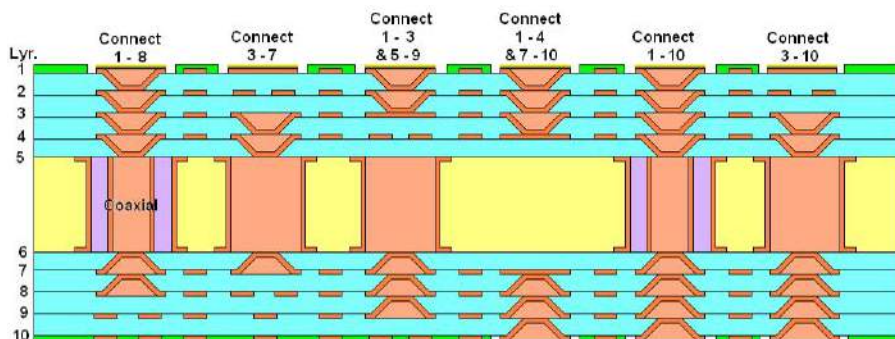
Na **rysunku 24** można także zobaczyć schemat płytki typu IV wg IPC-2226. W tej konfiguracji (spotykanej w praktyce nieporównanie rzadziej niż typy I...III) wewnątrz płytki stanowi tzw. rdzeń pasywny – niebiorący udziału w budowie obwodów, a jedynie pełniący funkcję radiatora, kompensatora rozszerzalności termicznej pozostałych warstw czy też dodatkowego ekranu.

Typ V wg IPC-2226 to już tzw. obwody bezrdzeniowe (**rysunek 25**). Wszystkie warstwy są wykonywane parami, z mikroprzelotkami łączącymi odpowiednie punkty obwodu, leżące po obydwu stronach każdego z dielektryków. Na koniec wszystkie arkusze są laminowane w ramach pojedynczego procesu (co odróżnia go od laminacji sekwencyjnej, stosowanej w typach II i III), zaś ewentualne połączenia pomiędzy warstwami leżącymi na sąsiadujących ze sobą arkuszach są wykonywane z użyciem... pasty przewodzącej prąd elektryczny.

Najbardziej złożonym typem struktury płyt HDI jest typ VI, określany mianem ELIC (od ang. *Every Layer InterConnect* bądź *any-layer HDI*). Jak sama nazwa wskazuje, możliwe jest połączenie ze sobą dowolnych zestawów sąsiadujących ze sobą warstw, gdyż wszystkie przejścia są wykonywane za pomocą mikroprzelotek piętrowych. Jak nietrudno zauważyć, sama definicja tej struktury stoi w opozycji do opisanego wcześniej zalecenia, aby unikać piętrowych zestawów zawierających więcej niż dwie mikroprzelotki (standard IPC-2226 powstał przed opracowaniem ww. zalecenia). Warto dodać, że nieliczni producenci oferują niebywale zaawansowaną technologię, określaną czasem mianem *coaxial via* lub *via-in-via*, polegającą na budowie przelotki... wewnątrz większej przelotki, wypełnionej



Rysunek 25. Budowa PCB HDI typu V wg IPC-2226 (<https://t.ly/jrwBj>)



Rysunek 26. Przykładowy przekrój 10-warstwowej płytki HDI w technologii *any-layer* (ELIC – *Every Layer InterConnect*) – <https://t.ly/bNJu7>



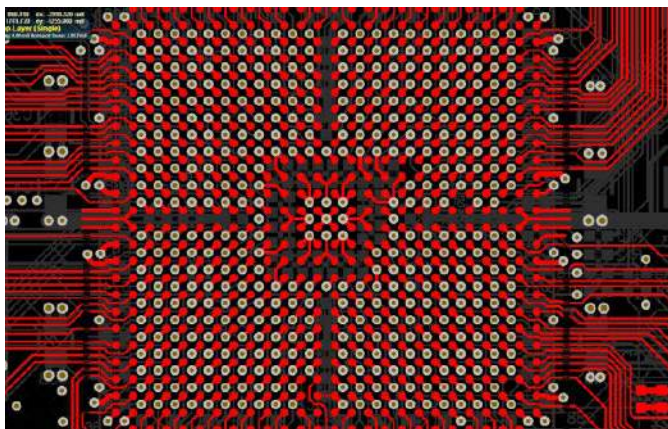
Rysunek 27. Proces wytwarzania przelotek w procesie CDD (*Controlled Depth Drilling*) – <https://t.ly/NC-cW>

uprzednio dielektrykiem. Przewiercenie utwardzonego izolatora i ponowna metalizacja umożliwia zatem „przepuszczenie” koncentrycznej przelotki w świetle innej, o nieco większej średnicy. Taką skomplikowaną strukturę pokazano schematycznie na **rysunku 26**.

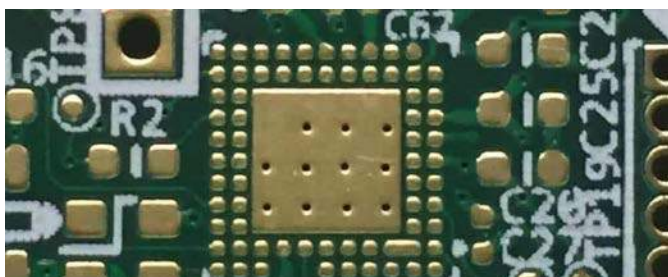
Omawiając złożoną tematykę przelotek, nie sposób nie wspomnieć o jeszcze jednej technice, określanej mianem CDD (ang. *Controlled Depth Drilling*) lub *back-drilled via*. Celem tej metody jest ograniczenie niepożądanej długiej części przelotki PTH, która łączy tylko kilka warstw położonych po jednej stronie PCB. Przelotka jest zatem wykonywana tak, jak zwykły otwór metalizowany na całej grubości płytki, po czym precyzyjnie ustawiona wiertarka numeryczna... usuwa niepotrzebną część metalizacji za pomocą wiertła (o średnicy nieznacznie większej niż średnica otworu bazowego), wprowadzonego na ściśle określoną głębokość. Schemat procesu pokazano na **rysunku 27**. A jaki sens ma rozwiercanie utworzonej wcześniej metalizacji? Po pierwsze, poprawia się w ten sposób integralność sygnałów, gdyż odbicia z używanej części przelotki pogarszają warunki transmisji. Po drugie, „wisząca” część metalizacji pełni funkcję miniaturowej anteny, która może nie tylko zwiększać emisję zakłóceń RFI, ale także podwyższać podatność czułego obwodu na przesłuchy z innych źródeł EMI.

Zagadnienia projektowania obwodów dla układów BGA

Układy w obudowach BGA są nierozdzielnie związane z technologią HDI, umożliwiają bowiem znaczną miniaturyzację PCB przy zachowaniu dużej liczby wyprowadzeń (często znacznie przekraczającej 100). Im większa macierz padów, tym więcej warstw może być potrzebnych do wyprowadzenia wszystkich niezbędnych linii sygnałowych i zasilających, zaś wraz ze zmniejszaniem rastra



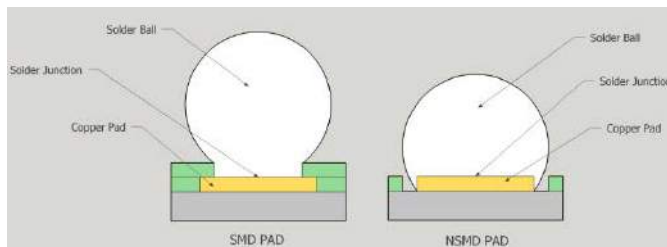
Rysunek 28. Fanout układu BGA z użyciem połączeń typu dog-bone (https://t.ly/6U_HJ)



Fotografia 7. Widok footprintu wykonanego z zastosowaniem technologii *via-in-pad* (<https://t.ly/YA18N>)

wyprowadzeń coraz trudniejsze może być zachowanie niezbędnych odstępów minimalnych oraz szerokości ścieżek, możliwych do wykonania przez wybranego producenta PCB. W przypadku układów o rastrze powyżej 0,5 mm dobrym i stosunkowo prostym rozwiązaniem (przynajmniej z technologicznego punktu widzenia) jest zastosowanie layoutu określanego mianem *dog-bone*. Polega on na połączeniu każdego z padów z niewielką przelotką, umieszczoną dokładnie pośrodku wolnej przestrzeni między sąsiadującymi padami – taka konstrukcja przypomina kształtem psią kość, stąd też charakterystyczna nazwa tej metody wyprowadzania pinów (tzw. *fanout*). Przykład można zobaczyć na **rysunku 28** – warto dodać, że ogromną pomocą dla projektantów są funkcje programów EDA, zapewniające automatyczny *fanout* padów (taką opcję można znaleźć m.in. w środowisku Altium Designer).

Drugą metodą na wyprowadzenie padów BGA do warstw leżących głębiej w stosie PCB jest zastosowanie przelotek w padach (*via-in-pad*), rzecz jasna wykonanych z dokładną planaryzacją (**fotografia 7**). Choć to rozwiązanie jest znacznie bardziej złożone technologicznie (wymaga bowiem najbardziej rozbudowanego procesu z użyciem przelotek



Rysunek 30. Porównanie typów padów BGA: SMD (po lewej) i NSMD (po prawej) – <https://t.ly/yGx49>



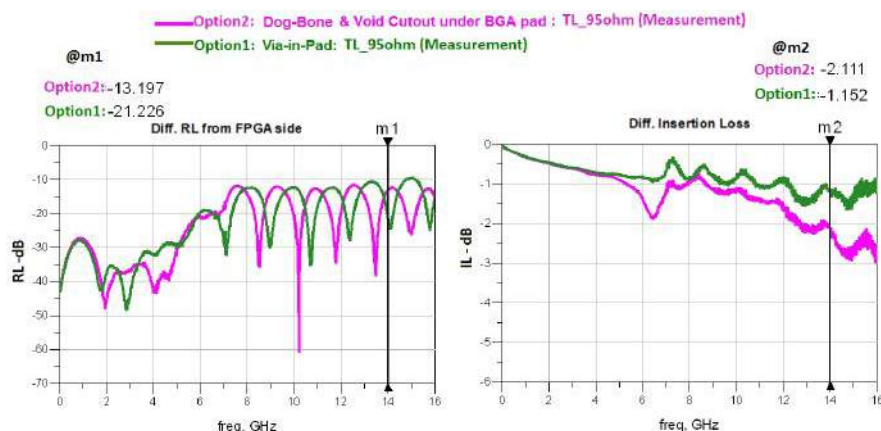
Fotografia 8. Przykładowe obwody typu *rigid-flex* (<https://t.ly/5Kho3>)

typu VII), to zapewnia szereg zalet: umożliwia pracę z rastrem BGA nawet znacznie poniżej 0,5 mm, znacząco poprawia odprowadzanie ciepła, podwyższa rating napięciowy (umożliwia bowiem pełne wykorzystanie dostępnych odstępów izolacyjnych), a jednocześnie – co też ważne – poprawia osiągi w zakresie integralności sygnałów. Choć na pierwszy rzut oka może się to wydawać mało prawdopodobne, nawet tak króciutki odcinek ścieżki w topologii *fanout* typu *dog-bone* (mający długość poniżej 1 mm) stanowi już dodatkową indukcyjność, której efekty są w pełni mierzalne. Za dowód niech posłużą wykresy z pomiarów porównawczych, dostępne w dokumentacji FPGA marki Intel – badanie strat odbiciowych i wtrąceńowych, przeprowadzone dla fanoutów *via-in-pad* oraz *dog-bone*, jasno wskazują wyższość tych pierwszych w zakresie ultraszybkich sygnałów cyfrowych i analogowych (zwłaszcza w paśmie powyżej 10 GHz – patrz **rysunek 29**).

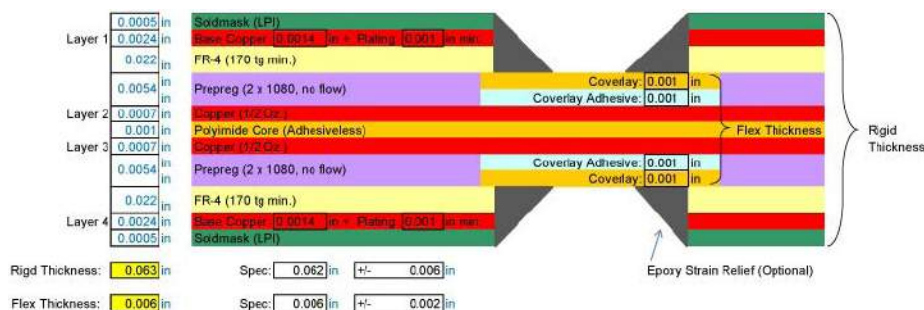
W przypadku padów BGA należy zwrócić uwagę na rodzaj zastosowanej techniki ograniczania powierzchni luty (**rysunek 30**). Dostępne są dwie metody: SMD (ang. *Solder Mask Defined*) oraz NSMD (ang. *Non-Solder Mask Defined*). W pierwszym przypadku średnica padu jest nieco większa, zaś ogranicza ją odpowiednio ukształtowane wycięcie w soldermasce, w drugim zaś powierzchnię luty definiuje faktyczny rozmiar padu, podczas gdy otwarcie soldermaski jest nieco szersze i w formie wąskiego pierścienia odsłania częściowo samo podłoże (prepreg lub rdzeń). W większości przypadków preferowaną metodą jest użycie padów NSMD, z uwagi na nieco prostszy proces produkcyjny (niewymagający wyśrubowanej precyzji nakładania soldermaski) oraz lepszy kontakt lutowni z powierzchnią padu.

PCB sztywno-giętkie

Na koniec zahaczmy jeszcze krótko o kolejny, niezwykle złożony temat obwodów sztywno-giętkich (*rigid-flex* – **fotografia 8**). Oprócz prostych płytek, pełniących funkcję



Rysunek 29. Porównanie strat odbiciowych i wtrąceńowych dla połączeń typu *dog-bone* oraz *via-in-pad* (<https://t.ly/plJrF>)



Rysunek 31. Budowa stosu prostego obwodu sztywno-giętkiego z czterema warstwami w części sztywnej i dwiema w części elastycznej. Na schemacie zaznaczono umocnienia, wykonane za pomocą żywicy epoksydowej w miejscach, w których obwód elastyczny opuszcza stos części sztywnej (<https://t.ly/oshXL>)



Rysunek 32. Budowa stosu rozbudowanego obwodu rigid-flex z 12 warstwami w części sztywnej i czterema w części elastycznej – kolorem szarym zaznaczono slot powietrzny pomiędzy dwoma pasmami laminatu elastycznego (<https://t.ly/QN1Ku>)

rozmaitych adapterów bądź fragmentów wewnętrznego okablowania urządzeń, spotkać można także bardzo rozbudowane konstrukcje, które zarówno w części sztywnej, jak i elastycznej, mają strukturę wielowarstwową. Przykład takiego rozwiązania pokazano na **rysunku 31** – dwuwarstwowa część giętka łączy dwie czterowarstwowe płytki, oparte na prepregach typu *no-flow* oraz usztywnieniach w postaci symetrycznie rozmieszczonych po obu stronach arkuszy FR4. Warto zwrócić uwagę na dwie istotne cechy: po pierwsze, konstrukcja stosu zmienia się w momencie przejścia z części sztywnej na giętką – zamiast prepregu występuje tutaj specjalna warstwa klejąca (określana mianem *bondply*), łącząca miedź z osłoną, wykonaną np. na bazie poliimidu. Z tego samego materiału jest także wykonany rdzeń części elastycznej, który z oczywistych przyczyn przechodzi przez stos na całej długości zestawienia. Po drugie, wspomniany interfejs stanowi punkt o krytycznym znaczeniu dla niezawodności połączeń, więc dobrym rozwiązaniem może być zastosowanie dodatkowej „odgiętki” w postaci cienkiej warstwy żywicy epoksydowej, wypełniającej narożniki utworzone przez pionowe ścianki części sztywnych oraz pokrycie części elastycznej.

Znacznie bardziej rozbudowany obwód zaprezentowano na **rysunku 32** – w tym przypadku części sztywne mają już po 12 warstw, zaś część elastyczna składa się z dwóch „taśm” dwuwarstwowch, przedzielonych cienką warstwą powietrza. Takie rozwiązanie zapewnia swego rodzaju „odciążenie mechaniczne”, ale zarazem wymagają bardziej złożoną konstrukcję całości.

Oprócz konstrukcji stosu, podczas projektowania obwodów wielowarstwowych typu *rigid-flex* należy oczywiście wziąć pod uwagę

szereg innych aspektów DFM. Odpowiedni dobór materiałów (w tym specjalnych prepregów, zdolnych do trwałego łączenia podłoża poliimidowych – zwykle kaptonu lub analogicznych materiałów – z miedzią), właściwe prowadzenie ścieżek (przy zachowaniu „łagodnej” geometrii i w miarę możliwości równoległej struktury głównych ścieżek) czy wreszcie dokładnie przemyślane umieszczenie ewentualnych komponentów, lutowanych na części elastycznej (z uwzględnieniem odpowiednich promieni gięcia w docelowym miejscu instalacji) to tylko wybrane spośród czynników, których uwzględnienie jest konieczne do osiągnięcia wymaganej niezawodności oraz akceptowalnych kosztów produkcji.

Podsumowanie

W artykule poruszyliśmy najważniejsze zagadnienia, związane z projektowaniem obwodów wielowarstwowych. Technologia nieustannie zmierza w kierunku coraz większej miniaturyzacji i wciąż przełamuje kolejne granice gęstości upakowania elementów na zaskakująco kompaktowych płytkach. Dziś mikroprzelotki o rozmiarach rzędu ułamka milimetra nie stanowią problemu dla tysięcy producentów PCB na całym świecie, a coraz więcej z nich oferuje już możliwości wytwarzania płyt, określanych mianem Ultra-HDI. Świat elektroniki (zwłaszcza konsumenckiej, np. smartfonów) powoli przestawia się na technologie semi-addtywne (SAP i mSAP), związane z wykonywaniem niewiarygodnie precyzyjnych obwodów drukowanych, których produkcja z użyciem standardowych procesów technologicznych byłaby już całkowicie niemożliwa. Ta tematyka wykracza jednak poza ramy niniejszego przeglądu fundamentalnych zagadnień techniki płyt wielowarstwowych – w niedalekiej przyszłości z pewnością powrócimy jednak do naszych Czytelników z opisami zaawansowanych technologii, stosowanych w supernowoczesnych płytkach drukowanych.

inż. Przemysław Musz, EP

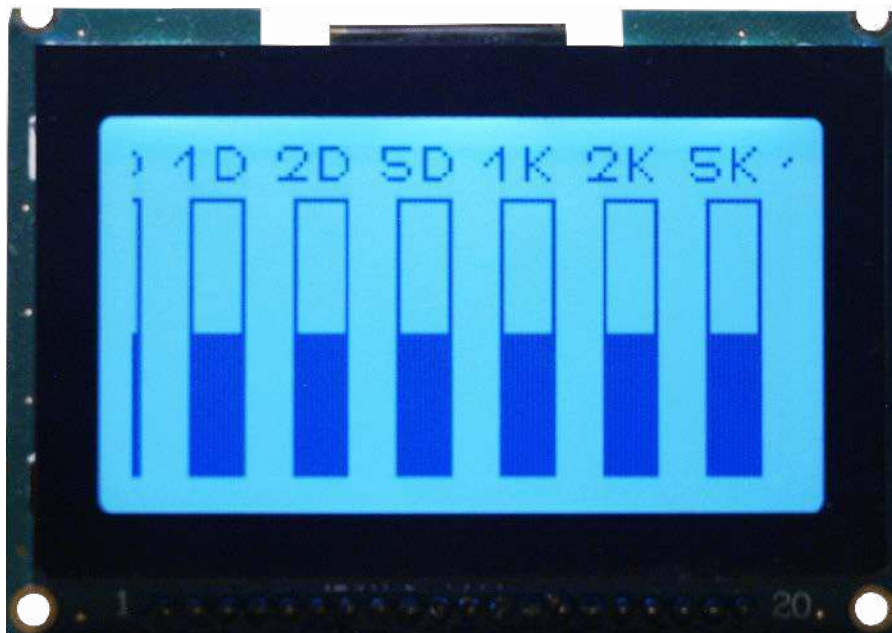
REKLAMA

www.feryster.pl

FERYSTER

Regulator barwy dźwięku na bazie płytki ewaluacyjnej LPCxpresso55S28

W artykule zaprezentuję ciekawą zastosowanie popularnego modułu wykonanego na bazie mikrokontrolera firmy NXP z rodziny LPC55Sxx – korektor barwy tonu. Zastosowana płytka ewaluacyjna – LPCxpresso55S28 – jest gotowym do pracy zestawem ze zintegrowanym programatorem i debuggerem oraz pełnym torem audio na bazie kodeka WM8904 firmy Cirrus Electronic. Sam mikrokontroler LPC55S28 jest wyposażony w interfejsy I²C oraz I²S przeznaczone do komunikacji z kodekiem oraz do przesyłania cyfrowego strumienia audio. Do artykułu dołączony jest pełny projekt programu.



Działanie urządzenia polega na przetworzeniu sygnału analogowego na cyfrowy poprzez przetwornik A/D zawarty w kodeku WM8904, a następnie rozdzieleniu go w zespolone pasmowych filtrów cyfrowych o różnych częstotliwościach. W każdym kanale częstotliwościowym można dokonać regulacji amplitudy sygnału. Następnie wartości

ze wszystkich kanałów są sumowane i podawane na przetwornik D/A kodeka.

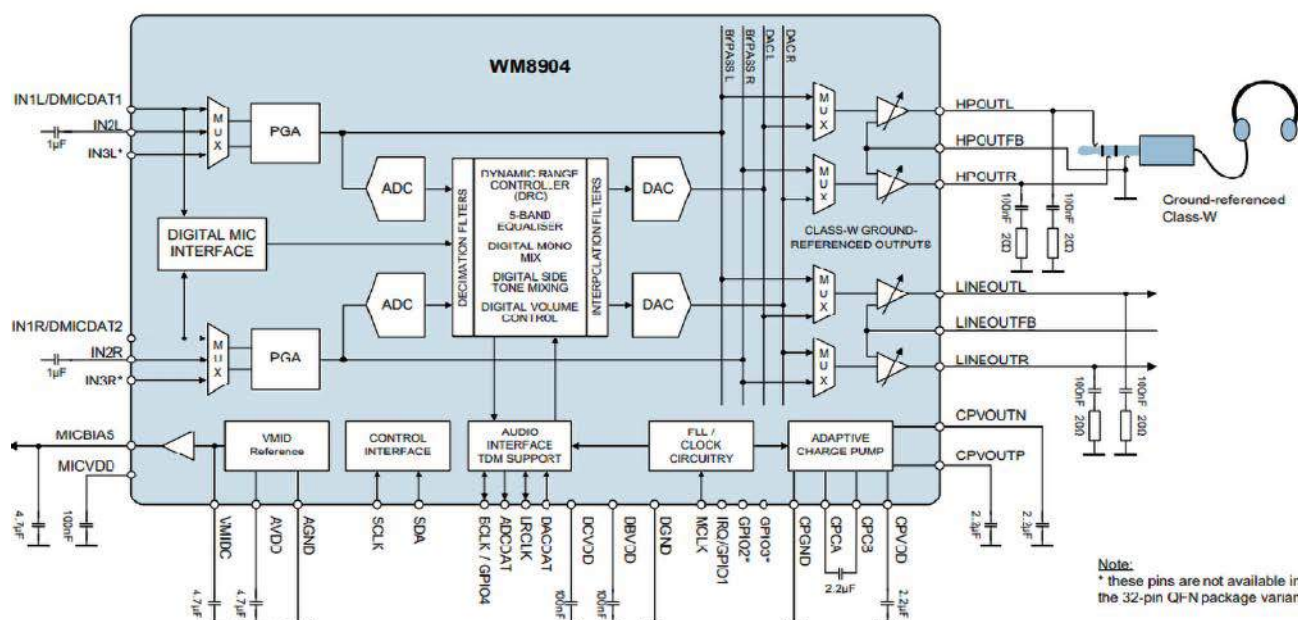
Opis programu komputerowego

Cały program składa się z dwóch części – procedur przetwarzania sygnału, w skład których wchodzi również komunikacja z kodekiem oraz

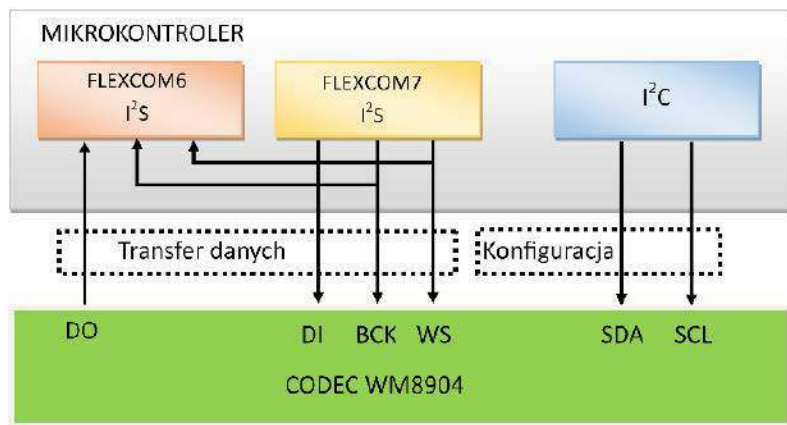
z interfejsu użytkownika, który jest odpowiedzialny za ustawianie parametrów urządzenia i obsługuje przyciski sprzętowe, enkoder oraz steruje wyświetlaczem graficznym.

Sterowanie kodekiem

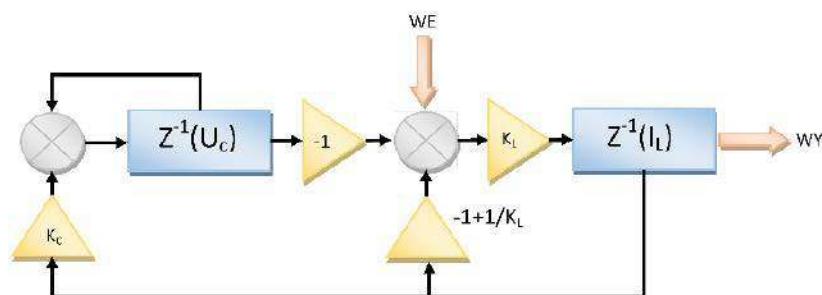
Sterowanie kodekiem WM8904 odbywa się poprzez dwa interfejsy: I²C oraz I²S.



Rysunek 1. Schemat blokowy kodeka WM8904



Rysunek 2. Realizacja transmisji pomiędzy MPU a kodekiem



Rysunek 3. Struktura filtra pasmowego realizowanego w programie

Magistrala I²C służy do konfiguracji układu. Proces ten polega na programowaniu wewnętrznych rejestrów procesora dźwięku. Można w ten sposób zaprogramować m.in.:

- wzmocnienie wewnętrznego wzmacniacza PGA,
- źródła sygnału wejściowego i wyjściowego,
- częstotliwość próbkowania, itp.

Interfejs I²S służy do przesyłania próbek dźwiękowych. W naszym programie zastosujemy gotową bibliotekę dźwiękową służącą do konfiguracji kodeka. Obsługa transmisji danych odbywa się poprzez dwa interfejsy I²S. Zastosowany mikrokontroler dostarcza system dziewięciu uniwersalnych bloków komunikacji szeregowej nazwanych FLEXCOM, każdy z nich można skonfigurować jako jeden z następujących interfejsów do transmisji danych:

- UART/USART
- SPI,
- I²C,
- I²S.

W naszym układzie użyjemy bloków FLEXCOM6 i FLEXCOM7, ponieważ są one elektrycznie podłączone do kodeka WM8904. Interfejs I²S zaimplementowany w module FLEXCOM6 jest używany jako wejście sygnału akustycznego, natomiast transmisja realizowana przez FLEXCOM7 wysyła dane do przetwornika DAC, czyli pełni funkcję wyjścia. Ponieważ sygnały sterujące transmisją (BCK i WS) są takie same dla obydwu interfejsów, należy dokonać ich połączenia. Służy do tego moduł SYSCTL.

Ponadto program ma funkcję regulacji sygnału wyjściowego oraz ustawienia typu wejścia: wejście mikrofonowe lub wejście liniowe.

Rodzaj wejścia odbywa się poprzez zmianę wzmocnienia wewnętrznego wzmacniacza PGA (*Programmable Gain Amplifier*). Obydwa te parametry można zmieniać poprzez ustawianie wartości w rejestrach kodeka.

Do działania układu WM8904 niezbędne jest dostarczenie sygnału taktującego na wejście MCLK. Sygnał ten wytwarzany jest przez mikrokontroler i wyprowadzony na odpowiednio skonfigurowane wyjście MCLK. Częstotliwość sygnału wynosi 24576000 Hz. Jest to zalecana wartość częstotliwości, która umożliwia wytworzenie wymaganych przebiegów sterującym przy pracy dla standardowych prędkości próbkowania. Na **rysunku 1** pokazano strukturę wewnętrzną kodeka, natomiast **rysunek 2** obrazuje transmisję pomiędzy kodekiem a mikrokontrolerem.

Opis algorytmu DSP

Cyfrowy strumień danych audio jest poddawany cyfrowej filtracji. Jako filtry zastosowałem cyfrowe imitacje szeregowych obwodów rezonansowych. Do ich zaprojektowania zastosowałem następujące zależności:

$$u_c = \frac{1}{C} \int i(t) \cdot dt$$

$$i_L = \frac{1}{L} \int u_L(t) \cdot dt$$

$$u_{we} = u_R + u_C + u_L$$

gdzie:

- U_R – napięcie na rezystorze,

Listing 1. Fragment kodu odpowiedzialny za przetwarzanie sygnału akustycznego

```
#define Filtr(f,q,x,y){\
    static float uc = 0;\\
    static float IL = 0;\\
    uc += IL * (float)\\
        ((float)2\\
        * STALA_PI / FP\\
        * (float)(q)\\
        * (float)(f));\\
    IL += (x - uc - IL)\\
        * (float)((float)2\\
        * STALA_PI / FP\\
        * (float)(f) / (float)(q));\\
    y = IL;}\n\nvoid __attribute__((section("RamFunction")))\nPrzetwarzanieGlowne() {\n\n    volatile register float kl = 0, kp = 0, wy, wl, wp;\\n\n    wl = (float)WartoscWE_L/(float)0x7FFF;\\n\n    wp = (float)WartoscWE_P/(float)0x7FFF;\\n\n\n    Filtr(21.5,1,wl,wy);\\n\n    kl+=wy*Poz.A20L;\\n\n    Filtr(21,1,wp,wy);\\n\n    kp+=wy*Poz.A20P;\\n\n\n    Filtr(46.4,1,wl,wy);\\n\n    kl+=wy*Poz.A50L;\\n\n    Filtr(46.4,1,wp,wy);\\n\n    kp+=wy*Poz.A50P;\\n\n\n    Filtr(100,1,wl,wy);\\n\n    kl+=wy*Poz.A100L;\\n\n    Filtr(100,1,wp,wy);\\n\n    kp+=wy*Poz.A100P;\\n\n\n    Filtr(215,1,wl,wy);\\n\n    kl+=wy*Poz.A200L;\\n\n    Filtr(215,1,wp,wy);\\n\n    kp+=wy*Poz.A200P;\\n\n\n    Filtr(464,1,wl,wy);\\n\n    kl+=wy*Poz.A50L;\\n\n    Filtr(464,1,wp,wy);\\n\n    kp+=wy*Poz.A500P;\\n\n\n    Filtr(1000,1,wl,wy);\\n\n    kl+=wy*Poz.A1KL;\\n\n    Filtr(1000,1,wp,wy);\\n\n    kp+=wy*Poz.A1KP;\\n\n\n    Filtr(2150,1,wl,wy);\\n\n    kl+=wy*Poz.A2KL;\\n\n    Filtr(2150,1,wp,wy);\\n\n    kp+=wy*Poz.A2KP;\\n\n\n    Filtr(4640,1,wl,wy);\\n\n    kl+=wy*Poz.A5KL;\\n\n    Filtr(4640,1,wp,wy);\\n\n    kp+=wy*Poz.A5KP;\\n\n\n    Filtr(10000,1.5,wl,wy);\\n\n    kl+=wy*Poz.A10KL;\\n\n    Filtr(10000,1.5,wp,wy);\\n\n    kp+=wy*Poz.A10KP;\\n\n\n    WartoscWY_L = (int16_t)(kl*(float)0x7FFF);\\n\n    WartoscWY_P = (int16_t)(kp*(float)0x7FFF);\\n\n}
```

- U_C – napięcie na kondensatorze,
- U_L – napięcie na cewce.

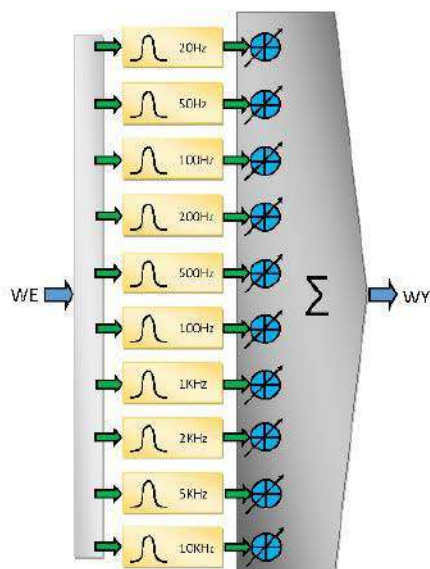
Wartości L i C można wyliczyć na podstawie częstotliwości rezonansowej oraz wymaganej dobroci filtru, natomiast wartość rezystancji najwygodniej przyjąć równą jeden. Struktura tego filtru zamieszczona jest na **rysunku 3**.

Fragment kodu odpowiedzialny za przetwarzanie sygnału akustycznego został pokazany na **listingu 1**. Aby użyć maksymalnej częstotliwości procesora procedura DSP, została umieszczona w pamięci operacyjnej.

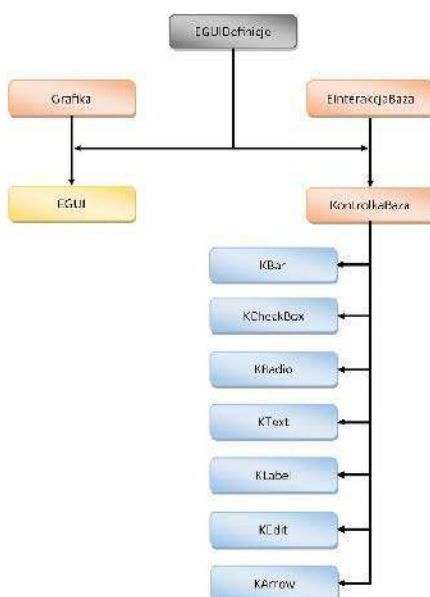
Na koniec wartości ze wszystkich filtrów są sumowane i podawane na przetwornik D/A kodeka, tak jak to pokazano na **rysunku 4**.

Interfejs GUI

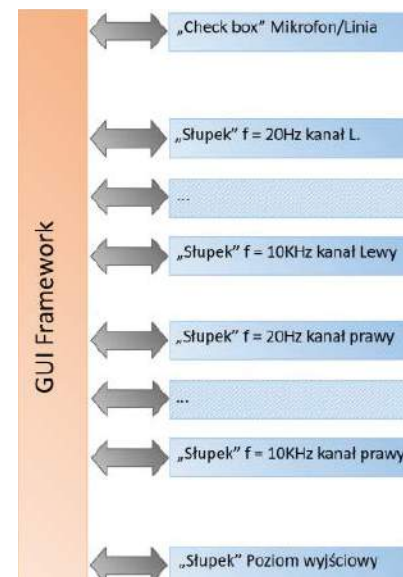
Do opracowania interfejsu użytkownika zastosowano zaprojektowaną przez autora bibliotekę, która współpracuje z wyświetlaczem graficznym LCD o rozdzielczości



Rysunek 4. Sposób tworzenia sygnału wyjściowego



Rysunek 5. Hierarchia klas interfejsu GUI



Rysunek 6. Schemat poglądowy struktury naszego interfejsu

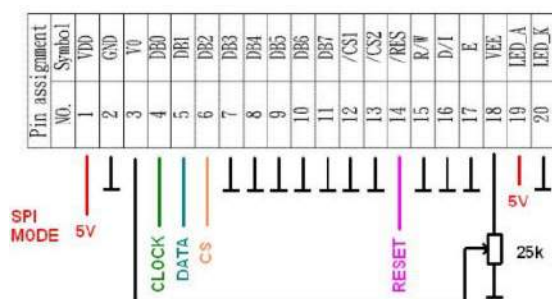
128×64 pikseli. Umożliwia ona utworzenie na ekranie 7 kontroltek:

- **Check box** – obiekt służy do wyboru dwóch stanów;
- **Radio Button** – ta kontrolka umożliwia wybór jednej z wielu opcji;
- **Text** – wyświetla tekst w ramce na ekranie i może być traktowana jako przycisk;
- **Label** – tylko wyświetla tekst;
- **Bar** – wyświetla słupkę o zmieniającej się wysokości. Kontrolka ta używana jest do ustawiania wartości;
- **Edit** – ten obiekt umożliwia edycję napisu;
- **Arrow** – wyświetla strzałkę na ekranie.

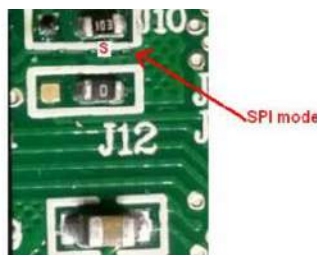
Główną klasą do obsługi interfejsu użytkownika jest *EGUI*. Umożliwia ona wyświetlanie treści na ekranie oraz odczyt wartości



Rysunek 7. Koncepcja działania interfejsu graficznego obsługującego wyświetlacz LCD



Rysunek 8. Schemat podłączenia wyświetlacza oraz sposób konfiguracji



z przycisków i enkodera. Również za pomocą jej funkcji składowych można do interfejsu dodawać okna kontrolne wywodzące się z klasy *EKontrolkaBaza*. Na **rysunku 5** pokazano hierarchię klas, natomiast na **rysunku 6** widzimy strukturę interfejsu GUI naszego programu.

Interfejs GUI korzysta z jeszcze jednej biblioteki graficznej opracowanej przez autora, służącej do wyświetlania różnych kształtów na ekranie wyświetlacza. Koncepcję działania tej biblioteki obrazuje **rysunek 7**.

Opis układu elektrycznego

Układ zbudowany jest na bazie gotowego modułu z mikrokontrolerem firmy NXP, ale do jego działania potrzebnych jest kilka elementów dodatkowych – wyświetlacz graficzny o rozdzielczości 128×64 piksele oraz prosty układ sterowania. Można dołączyć dwa przyciski, chociaż nie jest to konieczne, jeśli zdecydujemy się na sterowanie urządzeniem za pomocą przycisków znajdujących się na płycie ewaluacyjnej.

Zastosowany w naszym układzie wyświetlacz graficzny może pracować w dwóch trybach. Może być sterowany za pomocą transmisji równoległej lub szeregową, z zastosowaniem w transmisji SPI. W naszej aplikacji zastosujemy transmisję szeregową. Ponieważ wyświetlacz fabrycznie jest skonfigurowany do pracy przy transmisji równoległej, należy

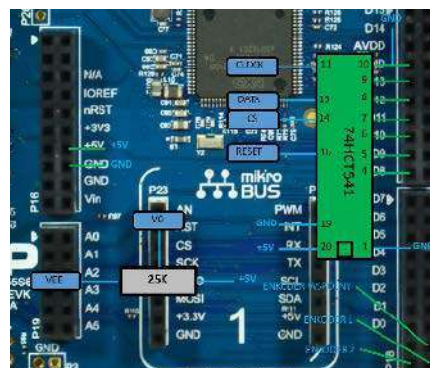
przełączyć zworę, która znajduje się na jego tylnej części (**rysunek 8**).

Dodatkowo, ponieważ układ logiczny wyświetlacza jest sterowany napięciami niezgodnymi z logiką standardu 3,3 V, czyli wartościami wyjściowymi mikrokontrolera (wysoki stan wejść wymaga napięcia wyższego niż 3,5 V), zastosowano bufor 74HCT541. Uproszczony schemat połączeń części elektrycznej pokazany jest na **rysunku 9**.

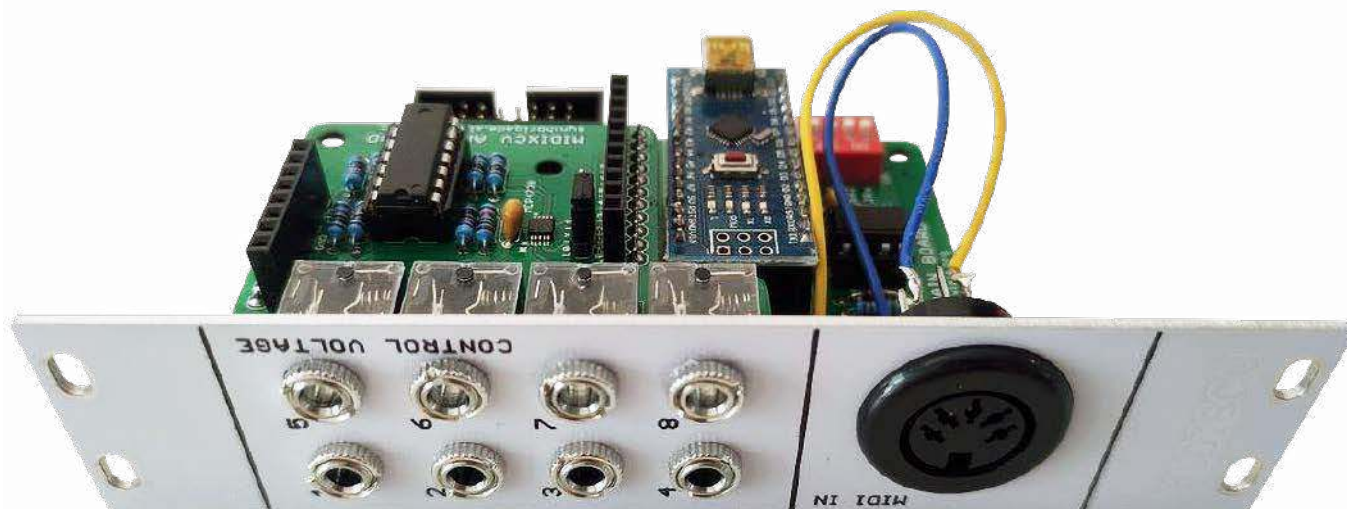
Elementy zastosowane w układzie

Do wizualizacji działania urządzenia zastosowano wyświetlacz monochromatyczny graficzny LCD o rozdzielczości 128×64 typu LCD-EG-128064H-FHW K/W-E6. Charakteryzuje się on niską ceną i jest dostępny w wielu sklepach internetowych m.in. Artronic czy Kamami. Jako bufor zastosowałem układ 74HCT541. Jest on szeroko dostępny, w niskiej cenie. Potencjometr do regulacji kontrastu może być dowolnym potencjometrem montażowym o rezystancji około 25 kΩ. Jako enkoder można zastosować dowolny 3-wyjściowy enkoder inkrementalny.

Tomasz Krogulski
krogul70@gmail.com



Rysunek 9. Schemat połączeń części elektrycznej



MIDIXCV

– konwerter MIDI do modułowego syntezyzatora analogowego

Modułowy syntezyzator audio to rodzaj syntezyzatora muzycznego zbudowanego z oddzielnych modułów pełniących różne funkcje. Moduły te można ze sobą łączyć, układać i konfigurować według potrzeb użytkownika, co daje ogromną elastyczność w tworzeniu różnych dźwięków i efektów dźwiękowych. Moduły te są zazwyczaj montowane na specjalnych panelach przednich, które mają złącza i potencjometry do łączenia i regulacji sygnałów.

Napięcia sterujące w modułowych syntezyzatorach audio są używane do kontrolowania różnych parametrów dźwięku, takich jak wysokość, głośność, czas trwania obwiedni i inne. Są one przekazywane między modułami w celu zmiany lub modyfikacji parametrów dźwięku. Na przykład napięcie sterujące z klawiatury może kontrolować wysokość dźwięku, a napięcie sterujące generatorem obwiedni może regulować kształt dźwięku w czasie.

Modułowe syntezyzatory audio pozwalają na tworzenie niemal nieograniczonej liczby brzmień i efektów dźwiękowych, co sprawia, że są popularne wśród muzyków, producentów dźwięku i entuzjastów muzyki elektronicznej. Dzięki napięciu sterującemu (CV) różnymi modułami można eksperymentować z dźwiękiem, tworzyć niestandardowe brzmienia i kontrolować niemal każdy aspekt dźwięku za pomocą elektronicznych modułów. Dlatego też napięcia sterujące są tak ważne w modułowych syntezyzatorach audio – to one pozwalają na sterowanie działaniem całego systemu. Może ono pochodzić z różnych źródeł – z jednej strony zadawane może być manualnie – potencjometrem-dzielnikiem napięcia, ale może pochodzić również z klawiatur sterujących, sekwencerów czy bardziej zaawansowanych urządzeń. W 2021 roku autor zaprezentowanej konstrukcji opublikował pierwszy projekt dotyczący konwertera MIDI na CV, który jest zdolny generować do czterech napięć kontrolnych i cztery sygnały cyfrowe na podstawie wiadomości MIDI. Zbudowany był wokół mikrokontrolera Arduino Nano i przetwornika DAC MCP4728. Przez lata pełnił funkcję mózgu w jego własnym parafonicznym syntezyzatorze modułowym, który własnoręcznie zbudował.

Ostatnio dodał do niego kilka funkcji do oprogramowania i niespodziewanie pojawiły się nowe pomysły i wystarczająca motywacja, aby wyjąć ten projekt z szuflady i zacząć coś nowego. Efektem tych prac jest wielokanałowy moduł do generowania napięć kontrolnych z ośmioma wyjściami analogowymi, z dostateczną rozdzielczością, aby obsługiwać dźwięki w różnych skalach muzycznych.

Ale to nie wszystko. Sprzęt, który został opracowany, może być rozszerzony do obsługi aż 32 (trzydziestu dwóch!) sterujących napięć analogowych i czterech sygnałów cyfrowych. Zakres oktaf, który był wcześniej ograniczony do czterech w przypadku konwertera MIDI na CV, został teraz rozszerzony do ośmiu.

Ponieważ oprogramowanie jest w pełni otwarte, kod może być modyfikowany w taki sposób, aby obsługiwać dowolne rodzaje przycho- dzących wiadomości MIDI (nuta włącz/wyłącz, numer nuty, prędkość, *pitch bend*, *aftertouch*, zmiana kontrolna, zmiana programu i tak dalej) i przekształcać je w odpowiednie napięcie za pomocą funkcji definiowanych przez użytkownika.

W artykule opiszemy moduł z ośmioma wyjściami i jego budowę oraz oprogramowanie, a także opiszemy, w jaki sposób łatwo dodać dodatkowe wyjścia analogowe. W dalszej części znajduje się opis sprzętu, który umożliwia uruchomienie takiego systemu, a także opis procedury programowania i kalibracji, bez potrzeby wykorzystania dodatkowych, specjalnych urządzeń.

Na stronie z projektem autor udostępnia wszystkie niezbędne pliki, aby można było wyprodukować potrzebne płytki, które składają się na ten moduł. Udostępnia też wszystkie niezbędne pliki do programowania modułu.

Potrzebne elementy

Oprócz płytek drukowanych potrzebnych jest trochę dodatkowych elementów do lutowania na nich. Są to głównie złącza i drobne elementy dyskretne. Poniżej wymieniono kluczowe elementy poszczególnych modułów.

Płytką główną

- 4 × złącza jack np. PJ324M (stereo lub mono),
- 1 × złącze MIDI DIN 5 (DS-5-01 do montażu na PCB lub DS-5-07A do montażu na panelu),

- 1 × 12-bitowy DAC MCP4728,
- 1 × transpator 6N138,
- 1 × wzmacniacz operacyjny TL074,
- 1 × złącze IDC 08x2,
- 1 × przełączniki DIP-4,
- diody 1N4004 (1 sztuka) i 1N4148 (1 sztuka),
- Oporniki: 5 × 1000 Ω, 1 × 220 Ω, 1 × 330 Ω, 2 × 4,7 kΩ, 8 × 10 kΩ,
- Kondensatory: 4 × 100 nF, 2 × 10 μF,
- Złącza 8 pin i 11 pin.

Płytki analogowa

- 4 × złącza jack np. PJ324M (stereo lub mono),
- 1 × 12-bitowy DAC MCP4728,
- 1 × wzmacniacz operacyjny TL074,
- Oporniki: 4 × 1000 Ω, 8 × 10 kΩ,
- Kondensator 100 nF,
- Złącza i wtyki 8 pin i 11 pin,
- Złącze 3 pin.

Płytki cyfrowa

- 4 × złącza jack np. PJ324M (stereo lub mono),
- 1 × wzmacniacz operacyjny TL074,
- Oporniki: 4 × 1000 Ω, 8 × 10 kΩ,
- Kondensator 100 nF,
- Złącza i wtyki 8 pin i 11 pin (męskie i żeńskie).

Moduł kalibracyjny

- 4 przyciski,
- Złącza (męskie) 8 pin i 11 pin.

Hardware

Projekt składa się z trzech płytek drukowanych:

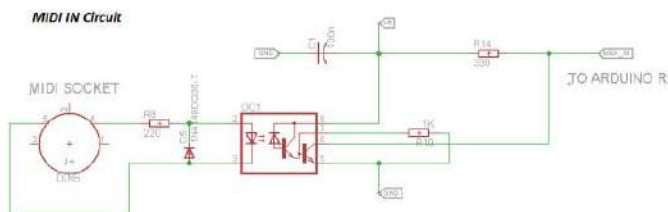
- Płytki główna
- Płytki analogowa
- Płytki cyfrowa

Dodatkowym elementem jest panel przedni, który umożliwia montaż systemu w modułowym syntezatorze. Autor przygotował do tego specjalny panel, do którego dokumentacja jest udostępniona razem z dokumentacją pozostałych płytek drukowanych w tym systemie.

Płytki główna

Ta płytka stanowi serce projektu. Zawiera moduł mikrokontrolera (Arduino NANO 328p), układ wejścia interfejsu MIDI, sekcję filtracji zasilania, cyfrowy stopień przetwarzania sygnału na analogowy oraz stopień wzmocnienia i ochrony wyjść. Płytki ma cztery wyjścia CV. Jeśli nasz projekt nie wymaga więcej, to możemy poprzestać na montażu tylko tego PCB.

Złącze MIDI (DIN-5) można bezpośrednio przylutować do tej płytki, ale spowoduje to, że zaprojektowany przez autora panel przedni nie



Rysunek 1. Schemat wejścia MIDI w układzie

będzie pasował do PCB – autor używa MIDIXCV w ten sposób zamocowanego we własnej obudowie DIY (jest trochę wolnej przestrzeni na dole). Alternatywnie można zastosować zgodny z EuroRack panel przedni, ale konieczne jest użycie innego złącza MIDI (fotografia 1).

Na tej płytce znajduje się wejście MIDI, zrealizowane w klasyczny sposób na pojedynczym transpatorze. Schemat tego wejścia pokazano na rysunku 1. Układ jest zgodny z specyfikacjami MIDI Association.

Przychodzące wiadomości cyfrowe MIDI muszą zostać przekonwertowane na napięcia analogowe, aby mogły komunikować się w tym samym języku co moduły syntezatora. Funkcję takiego konwertera pełni mikrokontroler w module Arduino NANO328 i przetwornik cyfrowo-analogowy (DAC). Mikrokontroler odbiera wiadomości MIDI i przekształca je w zdefiniowane przez użytkownika informacje w postaci cyfrowej. Następnie są one przesyłane do przetworników DAC za pomocą protokołu komunikacji I²C.

Za konwersję cyfrowych sygnałów na napięcia analogowe odpowiada poczwórny przetwornik DAC MCP4728 o rozdzielczości 12 bitów. Rozdzielczość przetwornika DAC wynosi około 1 mV (z wewnętrznym napięciem odniesienia), a na ostatnim etapie toru sygnałowego za wzmacniaczem spada do 2 mV. To bardzo dobre rozwiązanie sprzętowe dla tego zastosowania, ale także ważna jest dostępność wielu różnych bibliotek do tego przetwornika, co znacząco zwiększa szanse na sukces projektu hobbystycznego (nie zapominajmy o tym).

Zupełnie odrzucono użycie wyjść PWM Arduino we współpracy z (tanymi i łatwymi do zaprojektowania) filtrami RC. Rozwiązanie takie jest wysoce niestabilne (potrzebuje PWM o częstotliwości co najmniej 1 kHz, aby uzyskać sensowną jakość sygnału, a Arduino NANO oferuje mniej niż połowę tej częstotliwości).

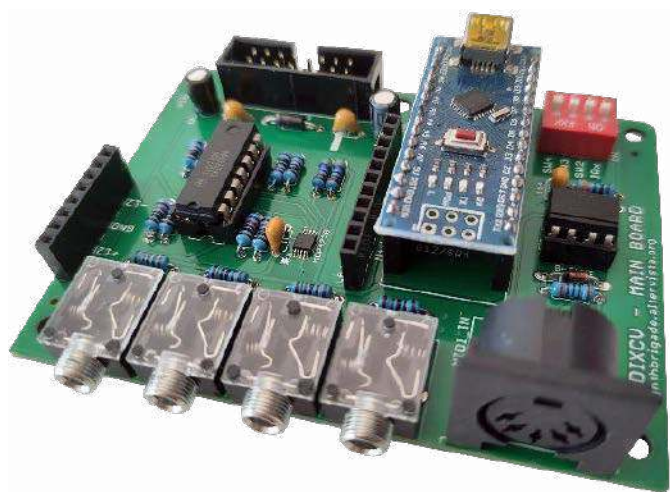
Przetwornik MCP4728 może być ustawiony programowo do korzystania z wewnętrznego napięcia odniesienia (stałe napięcie odniesienia 2,048 V) lub zewnętrznego (Vcc). Użycie wewnętrznego napięcia odniesienia ma zaletę stabilnego śledzenia napięcia, ale jest ograniczone do zakresu wyjściowego od 0 do 4,095 V. Oznacza to, że używając wewnętrznego napięcia odniesienia (bez wzmacniacza), mamy do dyspozycji zakres tylko czterech oktaf zamiast pięciu. Napięcie jest jednak bardziej powtarzalne i stabilniejsze przy użyciu wewnętrznego odniesienia. W tym projekcie DSC dla CV wysokości dźwięku używa wewnętrznego odniesienia napięcia DAC, a napięcia są podwajane przez końcowy wzmacniacz, co daje pełny zakres ośmiu oktafów.

Płytki analogowa

Napięcie analogowe – CV – to napięcie, które może przyjmować dowolną wartość w określonym zakresie napięcia (np. od 0 V do 10 V w przypadku typowych syntezatorów modułowych). Płytki analogowe (zauważcie, że jest tutaj liczba mnoga) służą do zwiększenia liczby wyjść analogowych.

Każda płytka analogowa dodaje cztery napięcia sterujące do całkowitej liczby wyjść modułu. Płytki analogowe oczywiście nie mogą działać same z siebie – są przeznaczone do montażu na głównej płytce. Im więcej wyjść analogowych potrzebuje użytkownik, tym więcej płytek analogowych należy ułożyć jedna na drugiej. Na przykład, jeśli w systemie potrzebne jest 8 wyjść CV, zostanie zmontowana i ułożona razem płytka bazowa i jedna płytka analogowa. Dla 12 wyjść trzeba zainstalować dwie płytki analogowe na płytce głównej.

W teorii można ułożyć na głównej płycie do siedmiu dodatkowych płytek analogowych, co oznacza ogólną liczbę trzydziestu dwóch



Fotografia 1. Płytki główna ze złączem na PCB do montażu w obudowie

napięć analogowych, ale autor informuje, że użycie więcej niż 12 wyjść nie zostało przetestowane. Konfiguracja 8 wyjść (płytką główną + jedna płytka analogowa) jest optymalną, zdaniem autora, konfiguracją, która była najczęściej testowana.

Każda z płytek analogowych ma kopię układu DAC z płytki głównej. Dzięki temu, że MCP4728 ma programowalny adres, można do pojedynczego interfejsu I²C w mikrokontrolerze podłączyć w teorii dowolną liczbę przetworników DAC.

Wyjścia DACa podłączone są do wzmacniaczy operacyjnych (TL074), które pozwalają na osiągnięcie napięcia 0..10 V (wyjście z DAC, osiąga maksymalnie 5 V). Wzmacniacz operacyjny zapewnia także zabezpieczenie przeciwzwarciowe wyjścia oraz zabezpieczenie przeciwprzepięciowe układu.

Na płytce znajdują się cztery przełączniki DIP. Jeden jest stałe podłączony do linii Rx Arduino i musi być używany do wgrywania nowego oprogramowania. Pozostałe trzy są bezpośrednio połączone z trzema cyfrowymi pinami Arduino i mogą być używane do konfiguracji modułu, np. do aktywacji dodatkowych funkcji i trybów zapisanych w kodzie źródłowym.

Płytką cyfrowa

Sygnały cyfrowe to sygnały używane do przesyłania wiadomości binarnych. Dwie wartości napięcia najczęściej reprezentują dwa stany – „0” i „1” – to 0 V i 5 V (lub 0 V i 3,3 V), ale w świecie modułowych syntezatorów napięcie stanu „1” może sięgać aż do 12 V. Mając to na uwadze, płytka cyfrowa nie była naprawdę konieczna, ponieważ napięcia analogowego można używać (w pewnym zakresie) do kontroli sygnałów cyfrowych. Jednakże płytka analogowa jest trudniejsza do zmontowania i droższa, więc autor zaprojektował osobną płytkę dla wyjść cyfrowych.

Płyta cyfrowa ma maksymalnie cztery niezależne wyjścia cyfrowe. Układanie więcej niż jednej płyty cyfrowej na płycie głównej nie zapewni więcej niż czterech niezależnych wyjść. Wyjścia te również są buforowane za pomocą wzmacniaczy operacyjnych.

Panel przedni

Ten specjalny panel został zaprojektowany do montażu płytki głównej i jednej dodatkowej płytki, dając łącznie osiem wyjść. Większej liczby niż jeden dodatkowy moduł nie da się zmieścić z obecnym panelem przednim, ale teoretycznie można byłoby zaprojektować nowe panele dla większej liczby modułów.

Panel eksponuje osiem wyjść i złącze MIDI IN użytkownikowi pozwala na montaż modułu w typowym syntezatorze modułowym np. w EuroRacku.

Firmware

W swoim syntezatorze autor przyjął najprostszą konfigurację parafoniczną, jaką można sobie wyobrazić, z czterema oscylatorami bezpośrednio podłączonymi do wyjść CV konwertera MIDI. Unikalność tej konfiguracji polega na tym, jak obsługiwane są głosy:

- Jeśli wciśnięto tylko jeden klawisz (sterowanie MIDI przesyła informacje o klawiszach), to wszystkie głosy mają tę samą wysokość dźwięku (podobnie jak w trybie unisono);
- Jeśli zostanie wciśnięty dodatkowy klawisz, to dwa z czterech głosów są przypisywane do nowego klawisza;
- Jeśli trzeci klawisz zostanie dodany, to dwa głosy są wygaszane i przypisywane do nowego klawisza;

Listing 1. Kod mikrokontrolera (fragmenty)

```
#include <MIDI.h>
#include <Wire.h>
#include "mcp4728.h"

// Inicjalizacja obiektu MCP4728, ID = 0
mcp4728 dac0 = mcp4728(0);
// Inicjalizacja obiektu MCP4728, ID = 1
mcp4728 dac1 = mcp4728(1);

#define MAX_VOICES 4
#define MAX_INT_V 97
#define MIDI_CHANNEL 1
#define MIDIOFFSET 24

#define GATE 0
#define VEL 1
#define KBTRACK 2
#define OPEN 2500
#define CLOSED 0
// Ograniczenie napięcia AUX do 5 V
#define LIMIT_CV

const byte LDAC0 = 2;
const byte DIP1 = 10;
const byte DIP2 = 11;
const byte DIP3 = 12;
bool DIP1State;
bool DIP2State;
bool DIP3State;

MIDI_CREATE_DEFAULT_INSTANCE();

int noteCount;
byte lowestNote;
byte highestNote;
int activeSlot;
int pitchbend = 0;
int noteOverflow;
int velVal;
int kbtVal;
int ATVoltage;

byte noteMem[MAX_VOICES];
bool busy[MAX_VOICES];

// Tabela napięć w mV/półton dla DAC #1 (powtarza się dla każdego DAC
const int cv0IntRef[MAX_INT_V] = {
0, 44, 88, 130, 172, 215, 257, 299, 342, 384, 426, 468,
508, 550, 592, 633, 675, 718, 759, 801, 843, 884, 926, 968,
1007, 1049, 1091, 1133, 1174, 1218, 1259, 1301, 1344, 1385, 1426, 1468,
1510, 1551, 1593, 1635, 1677, 1720, 1762, 1803, 1846, 1887, 1929, 1971,
2012, 2053, 2095, 2137, 2179, 2221, 2263, 2304, 2347, 2389, 2431, 2473,
2515, 2554, 2596, 2638, 2680, 2722, 2763, 2805, 2847, 2889, 2931, 2973,
3015, 3056, 3097, 3140, 3181, 3223, 3265, 3307, 3349, 3392, 3434, 3477,
3519, 3560, 3601, 3643, 3685, 3727, 3770, 3812, 3854, 3896, 3938, 3981,
4022
}; //4095 MAX

void setup(){
//DIP switche
pinMode(DIP1, INPUT_PULLUP);
pinMode(DIP2, INPUT_PULLUP);
pinMode(DIP3, INPUT_PULLUP);
DIP1State = digitalRead(DIP1);
DIP2State = digitalRead(DIP2);
DIP3State = digitalRead(DIP3);
//Inicjalizacja MIDI
MIDI.setHandleNoteOn(HandleNoteOn);
MIDI.setHandleNoteOff(HandleNoteOff);
MIDI.setHandlePitchBend(HandlePitchBend);
MIDI.setHandleAfterTouchChannel(HandleAfterTouch);
MIDI.begin(MIDI_CHANNEL);
//Inicjalizacja DAC
pinMode(LDAC0, OUTPUT);
digitalWrite(LDAC0, LOW);
dac0.begin(); // Inicjalizacja Interfejsu I2C
dac1.begin(); // Inicjalizacja Interfejsu I2C
dac0.setPowerDown(0, 0, 0, 0);
dac1.setPowerDown(0, 0, 0, 0);
// Wewnętrzne napięcie odniesienia
dac0.setVref(1, 1, 1);
dac1.setVref(1, 1, 1);
// Wzmocnienie dla napięcia odniesienia
// (0 = x1, 1 = x2)
dac0.setGain(1, 1, 1, 1);
dac1.setGain(1, 1, 1, 1);
// Wyzeroowanie wszystkich wyjść
dac0.analogWrite(0, 0, 0, 0);
dac1.analogWrite(0, 0, 0, 0);
// Inicjalizacja zliczania dźwięków
noteCount = 0;
}

void HandleNoteOn(byte channel, byte note, byte velocity) {
noteCount++;
note = note-MIDIOFFSET;
if(note < MAX_INT_V && note>=0){
switch (noteCount){
case 0: //do nothing
break;
// jeden dźwięk - wszystkie głosy takie same,
// analogicznie dalsze przypadki
case 1:
for (int a = 0; a < MAX_VOICES; a++){
noteMem[a] = note;
busy[a] = true;
}
#ifdef LIMIT_CV
velVal = velocity << 3;

```

- Jeśli wciśnięte zostaną cztery klawisze, wszystkie oscylatory są niezależnie ustawiane na różne wysokości dźwięku.

Syntezator działa w trybie „ubogiego unisono”, jak określa to autor, z zawsze aktywnymi oscylatorami audio. Za każdym razem, gdy naciśnięty zostanie nowy klawisz, sygnał bramki jest ponownie uruchamiany, a co za tym idzie, generowana jest kolejna obwiednia.

Podejście parafoniczne ma dwie zalety w porównaniu z konwencjonalnym:

- pozwala uniknąć multipleksera / VCA w celu wyciszenia nieużywanych głosów,
- pozwala używać pojedynczego sygnału bramki zamiast oddzielnych bramek, po jednej dla każdego VCA.
- dźwięk jest pełniejszy w przypadku mniejszej liczby głosów niż cztery.

Oprogramowanie obsługuje osiem niezależnych wyjść analogowych. Wyjścia są poukładane w następujący sposób:

- Cztery wyjścia główne płyty są używane do generowania napięć kontrolnych w zależności od otrzymanej wysokości dźwięku MIDI (format V/oktawę, co oznacza, że każdy półton dzieli napięcie 0,83 mV od kolejnego);
- Wyjścia z płytki analogowej:
 - #1 wysyła sygnały bramki (pojedyncza brama – ustawienie parafoniczne „ubogie unisono”),
 - #2 wysyła napięcie proporcjonalne do ostatniej prędkości dźwięku,
 - #3 wysyła napięcie proporcjonalne do pozycji nuty na klawiaturze,
 - #4 wysyła napięcie proporcjonalne do *aftertouch*.

Obecne oprogramowanie obsługuje również *pitchbend*, jednakże nie jest ono przekazywane do CV w domyślnej konfiguracji.

Wszystkie wyjścia głównej płytki mają własną tabelę CV. Oznacza to, że procedura kalibracji (patrz opisany poniżej etap kalibracji) musi być powtórzona cztery razy, po jednym razie dla każdego wyjścia. Analogowe napięcia dla dodatkowych CV są programowo ograniczone do około 5 V, ale można zwiększyć zakres do 8 V, komentując linię: `#define LIMIT_CV` w kodzie (co oznacza umieszczenie znaków „//” przed linią lub jej usunięcie).

Przełączniki DIP, w obecnej wersji firmware pozostawione, są nieużywane, z wyjątkiem DIP1. DIP1 jest połączony bezpośrednio z linią Rx Arduino i musi być używany do wgrywania nowego oprogramowania. Aby wgrać nowy program, należy przełączyć DIP1 w pozycję „OFF”. Po zakończeniu wgrywania należy przełączyć go ponownie w pozycję „ON”, aby odbierać wiadomości MIDI.

Istotne fragmenty firmware pokazane są na **listingu 1**, a kompletny kod programu znajduje się w repozytorium autora na GitHubie (link na końcu artykułu). Aby skompilować główny program, wymagane są następujące biblioteki zainstalowane w środowisku Arduino IDE:

- Biblioteka MIDI od FortySevenEffects,
- Biblioteka MCP4728 autorstwa Benoit Shillingsa.

Linki do pobrania i instalacji bibliotek znajdują się na końcu artykułu.

Należy zauważyć, że stworzony przez autora program jest tylko (w 100% działającym!) programem demonstracyjnym, który ma na celu pokazanie pewnych możliwości układu – takich jakich autor używa w swoim syntezatorze, ale MIDIXCV ma oczywiście o wiele szersze możliwości i firmware może zostać dostosowane do ich używania.

Kalibracja wyjścia

Chociaż w teorii napięcia do sterowania wysokością dźwięku można łatwo obliczyć, fizyka rzeczywistego świata sprawia, że sprawy stają się bardziej skomplikowane. Dodatkowo, ponieważ nie ma dwóch

Listing 1. Kod mikrokontrolera (fragmenty) – cd.

```

    kbtVal = note << 3;
  #else
    velVal = velocity << 4;
    kbtVal = note << 4;
  #endif
  dac1.analogWrite(OPEN, velVal, kbtVal, ATVoltage);
  lowestNote = note;
  break;
// ponad cztery dźwięki - przekroczona liczba głosów
default:
  #ifdef LIMIT_CV
    velVal = velocity << 3;
    kbtVal = note << 3;
  #else
    velVal = velocity << 4;
    kbtVal = note << 4;
  #endif
  dac1.analogWrite(CLOSED, velVal, kbtVal, ATVoltage);
  for (int n = 0; n < MAX_VOICES; n++){
    if (noteMem[n] == lowestNote){
      switch(n){
        case 0:
          noteMem[0] = note;
          busy[0] = true;
          break;
        case 1:
          noteMem[1] = note;
          busy[1] = true;
          break;
        case 2:
          noteMem[2] = note;
          busy[2] = true;
          break;
        case 3:
          noteMem[3] = note;
          busy[3] = true;
          break;
      }
    }
  }
  NoteLowest();
  noteCount = MAX_VOICES;
//GATE 1 OPEN to complete the retrigger routine
  dac1.analogWrite(GATE, OPEN);
  break;
} //switch close
}
dac0.analogWrite(cv0IntRef[noteMem[0]]
+ pitchbend, cv1IntRef[noteMem[1]]
+ pitchbend, cv2IntRef[noteMem[2]]
+ pitchbend, cv3IntRef[noteMem[3]]
+ pitchbend);
}

void HandlePitchBend(byte channel, int bend){
  pitchbend = bend>>6;
  dac0.analogWrite(cv0IntRef[noteMem[0]]
+ pitchbend, cv1IntRef[noteMem[1]]
+ pitchbend, cv2IntRef[noteMem[2]]
+ pitchbend, cv3IntRef[noteMem[3]]
+ pitchbend);
}

void HandleAfterTouch(byte channel, byte pressure){
  #ifdef LIMIT_CV
    ATVoltage = pressure << 3; // 255->2040
  #else
    ATVoltage = pressure << 4; // 255->4080
  #endif
  dac1.analogWrite(3, ATVoltage);
}

void loop(){
  MIDI.read();
}

```

identycznych komponentów elektronicznych, proces kalibracji jest obowiązkowy, aby układ działał precyzyjnie.

Kalibracja polega na precyzyjnym dostrojeniu wszystkich pojedynczych napięć wyjściowych. W głównym szkicu są one przypisane do liczb od zera do 4095 (12-bitowa). Celem tego procesu jest znalezienie liczby, która zapewni jak najdokładniejsze dopasowanie rzeczywistego napięcia do idealnego napięcia dla każdej pojedynczej półtonacji w obszarze obsługiwanym przez całe oktawy. Kalibracja wymaga dobrego, skalibrowanego woltomierza oraz komputera PC z zainstalowanym Arduino IDE. Wyniki są wyświetlane na żywo przez monitor szeregowy.

Aby skalibrować układ do mikrokontrolera, należy wgrać specjalny szkic. Obsługuje on 4 niezależne wyjścia, z tabelą napięć dla każdego z nich. Oznacza to, że trzeba dostroić 96×4 wartości.

Procedura kalibracji

Aby rozpocząć kalibrację modułu, należy załadować do mikrokontrolera szkic do kalibracji (*MIDIXCV_outs_calibration.ino*), dostępny w repozytorium na GitHubie. Następnie pozostawiając

moduł podłączony do komputera, uruchamiamy monitor szeregowy w Arduino IDE i podłączamy woltomierz do wyjścia analogowego, które chcemy skalibrować, a następnie:

1. Zwieramy do masy pole „N+” lub „N-” na głównej płycie, aby zwiększyć lub zmniejszyć numer głosi MIDI, który chcemy dostroić.
2. Zwieramy do masy pole „V+” lub „V-” na głównej płycie, aby zwiększyć lub zmniejszyć napięcie bieżącego dźwięku MIDI o jeden krok (około 2 mV).
3. Gdy wszystkie dźwięki zostaną dostrojone, naciskamy „N-”, jednocześnie trzymając wciśnięte „N+” (dźwięk +): zostanie wyświetlona na porcie szeregowym pełna tablica.
4. Należy skopiować i wkleić dane do szkicu *MIDIXCV.ino* (patrz listing 1), zastępując wartości w tablicy `cv0IntRef[MAX_INT_V]` dla kanału 0 i tak dalej.
5. Operacja jest powtarzana dla wszystkich czterech wyjść.
6. Finalnie wgrywamy zmodyfikowany *MIDIXCV.ino* z czterema zaktualizowanymi tablicami do Arduino.

Kalibrację należy przeprowadzać z obciążeniem – wejście sterujące modułu VCO to bardzo dobre rozwiązanie, a w razie potrzeby można użyć po prostu opornika 100 kΩ do masy (np. wartość obciążenia w module VCO autora, opartym na AS3340, wynosi 90 kΩ).

Ponieważ procedura kalibracji może być uciążliwa, autor przygotował prostą płytkę z przyciskami kalibracyjnymi w celu ułatwienia pracy. Jej użycie nie jest obowiązkowe, ale znacznie ułatwia pracę.

Domyślnie szkic kalibracji jest ustawiony do kalibracji wyjść urządzenia z identyfikatorem #0, ale można go łatwo dostosować do dowolnych innych możliwych urządzeń (ID# 1 do 7). Jest to unikalne ID dla każdego przetwornika DAC w systemie. W kolejnym kroku opisano sposób konfigurowania tych ID dla poszczególnych modułów.

Konfiguracja płytek analogowych

Aby nadać poszczególnym przetwornikom DAC w systemie unikalne ID, co upraszcza komunikację z mikrokontrolerem, konieczne jest przeprogramowanie ich z unikalnym adresem każdego z nich. Jak zawsze, tak i tutaj społeczność Arduino przychodzi z pomocą i umożliwia operacje takie jak ta, bez konieczności posiadania szerokiej wiedzy i wielkich umiejętności programistycznych. Autor zastosował nieco zmodyfikowany szkic do reprogramowania od Jana Knippera (patrz link na końcu artykułu), aby ułatwić sobie pracę. Wystarczy postępować zgodnie z poniższymi krokami:

1. Zainstaluj bibliotekę *SoftI2CMaster* w swoim środowisku Arduino IDE (link do repozytorium z biblioteką na końcu artykułu).
2. Uruchom szkic do reprogramowania (*MIDIXCV_DAC_address_prog.ino*) z repozytorium projektu.
3. Zmień globalną stałą `#NEW_ADDRESS` na wartość, którą chcesz nadać danemu DAC (jeśli masz tylko jedną płytkę analogową, ustaw ją na „1”, w przeciwnym razie użyj indywidualnych wartości dla każdej płytki analogowej).

4. Połóż na płycie głównej płytę analogową, którą zamierzasz skonfigurować.
5. Połącz pin „LX” i „L1” za pomocą przewodu połączeniowego (zworki).
6. Wgraj szkic i uruchom szkic.
7. Po wgraniu szkicu należy odłączyć zworkę i przywrócić połączenie na pinach „LX” i „L0”.

Można przetestować, ile układów DAC o unikalnych adresach jest zainstalowanych, uruchamiając program *I2C_Scanner*. W monitorze szeregowym program powinien wyświetlić wszystkie informacje – pokazane zostaną unikalne adresy DAC, jakie wygenerowane zostały dla każdego z przetworników DAC.

Podsumowanie

MIDIXCV to ciekawa konstrukcja – konwerter MIDI na napięcie sterujące do modułowego syntezyzatora analogowego. Powyższy artykuł prezentuje to rozwiązanie, które pozwala na przekształcenie sygnałów MIDI na napięcia kontrolne, które są stosowane w modułowych syntezyzatorach audio do regulacji parametrów dźwięku. Syntezyzatory tego rodzaju pozwalają na elastyczne tworzenie różnych brzmień i efektów dźwiękowych poprzez łączenie i konfigurowanie różnych modułów. Dodanie sterowania z pomocą MIDI tylko zwiększa ich możliwości.

W artykule opisano wszystkie aspekty, potrzebne do budowy tego rodzaju urządzenia – hardware układu, z szczególnym naciskiem na przetwornik cyfrowo-analogowy i tor wyjściowy oparty o wzmacniacz operacyjny oraz oprogramowanie i zasada działania w syntezyzatorze parafonicznym. Dodatkowo dostarczono oprogramowanie do konfiguracji modułów oraz kalibracji poszczególnych kanałów wyjść analogowych.

To, że prezentowane oprogramowanie jest otwarte, pozwala w dowolny sposób je modyfikować i dostosowywać do różnych potrzeb, obsługując różne rodzaje przychodzących wiadomości MIDI i różnego rodzaju algorytmy sterowania syntezyzatorami. Dzięki temu projektowi entuzjaści muzyki i syntezyzatorów modułowych DIY mogą tworzyć własne, niestandardowe brzmienia za pomocą modułu syntezyzatora audio.

Nikodem Czechowski, EP

Bibliografia:

1. <https://shorturl.at/bpC56>
2. <https://github.com/baritonomarchetto/MIDIXCV>
3. https://github.com/FortySevenEffects/arduino_midi_library
4. <https://github.com/BenoitSchillings/mcp4728>
5. https://github.com/jknipper/mcp4728_program_address
6. <https://github.com/TrippyLighting/SoftI2cMaster>

REKLAMA

Pobierz bezpłatnie multimedialne dodatki do tego wydania „Elektroniki Praktycznej” projekty, miniprojekty, materiały do artykułów i kursów oraz wiele innych!

**• Kupiłeś magazyn w Ulubionym Kiosku lub masz prenumeratę?
Multimedialne dodatki będą odblokowane automatycznie!**

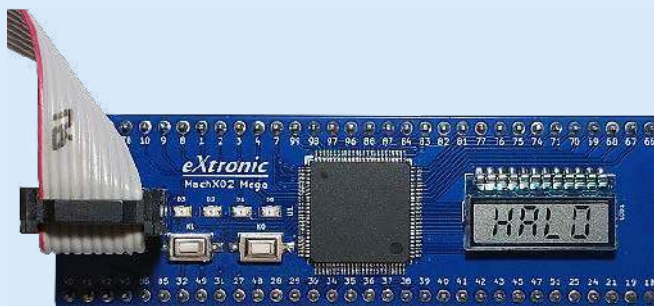
**• Zakupiłeś czasopismo u zewnętrznego dystrybutora?
Odblokuj bibliotekę multimedialną samodzielnie.**

Szczegóły na UlubionyKiosk.pl/media

Kurs FPGA Lattice (13)

Wyświetlacz LCD multipleksowany

Wyświetlacze LCD są wszędzie. Są tanie, energooszczędne i mogą wyświetlać najróżniejsze kształty. Jednak ich obsługa jest bardziej problematyczna niż wyświetlacza LED – multipleksacja jest skomplikowana i często wymaga aż czterech różnych napięć zasilających. Z tego powodu konieczne jest zastosowanie specjalizowanego sterownika LCD lub zbudowanie własnego.



Poprzednie odcinki znajdują się pod adresem:
<https://ulubionykiosk.pl/media>

Multipleksację wyświetlacza LED już znamy. Była omawiana w 9 odcinku kursu. Polega na tym, że każda cyfra uaktywniana jest za pomocą osobnej elektrody. Możliwe jest wyświetlanie tylko jednej cyfry jednocześnie, ale dzięki szybkiemu przełączaniu cyfr ludzkie oko widzi wyświetlacz tak, jakby wszystkie cyfry świeciły się jednocześnie. Elektrody odpowiadających sobie segmentów (oznaczonych od A do G oraz segment kropki P) są ze sobą połączone. W tak zorganizowanym wyświetlaczu mamy 8 elektrod segmentów i tyle elektrod wspólnych, ile jest cyfr w wyświetlaczu. Wszystkie elektrody sterowane są stanem niskim lub wysokim i nie ma żadnych stanów pośrednich.

Jak działa multipleksowany wyświetlacz LCD?

Multipleksacja wyświetlacza LCD niestety jest dużo trudniejsza. Na **rysunku 1** pokazano schemat połączeń wyświetlacza LCD-S401M16KR, który zastosowano na płycie MachXO2-Mega. Taki układ jest często stosowany w wyświetlaczach cyfrowych LCD wielu producentów, ale niestety nie jest standardem stosowanym przez wszystkich. Pomimo że wyświetlacz ma cztery cyfry i cztery elektrody wspólne COM, to elektrody wspólne nie aktywują pojedynczych cyfr, a grupy segmentów we wszystkich cyfrach. Aby zaczernić segment, musimy uaktywnić elektrodę COM i jednocześnie uaktywnić odpowiadającą mu elektrodę SEG. Na przykład aktywowanie elektrody COM3 sprawia, że możliwe staje się wyświetlenie segmentów A oraz F we wszystkich cyfrach, jeżeli aktywowane są także odpowiadające im elektrody SEG0...SEG7. Taki sposób sterowania sprawia, że moduł wyświetlacza multipleksowanego, jaki opracowaliśmy w 9 odcinku kursu, jest w tym przypadku zupełnie bezużyteczny.

W tym momencie musimy wprowadzić jeden z dwóch kluczowych parametrów wyświetlaczy LCD, jakich używają ich producenci – **duty**. Parametr ten mówi, przez jaki czas cyklu aktywna jest jedna elektroda wspólna (i jednocześnie ile jest elektrod wspólnych). Typowo stosuje się:

- **1 duty** – brak multipleksacji, 1 elektroda wspólna sterująca wszystkimi segmentami,
- **1/2 duty** – dwie elektrody wspólne,
- **1/3 duty** – trzy elektrody wspólne,
- **1/4 duty** – cztery elektrody wspólne,
- i tak dalej...

Sprawę dodatkowo komplikuje fakt, że elektrody większości wyświetlaczy LCD sterować trzeba nie tylko stanem wysokim czy niskim, ale także stanami pośrednimi. Tutaj pojawia się kolejny parametr, czyli **bias**, określający, jakimi napięciami należy sterować piny wyświetlacza. Możliwe są następujące opcje (gdzie 1 oznacza napięcie zasilania):

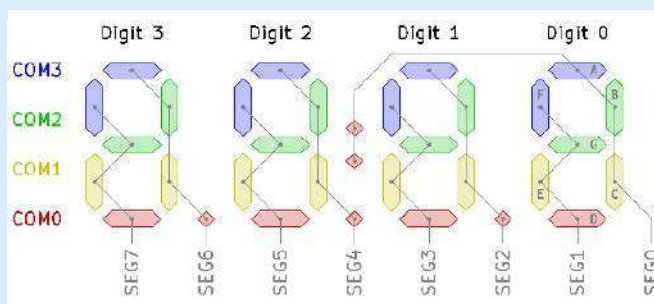
- Projekt w Diamond i pliki źródłowe możesz pobrać z serwera EP: <https://shorturl.at/cellR>
- Film demonstrujący działanie sterownika LCD z tego odcinka kursu: <https://youtu.be/6xA3YP8uHJs>
- Repozytorium autora kursu: <https://github.com/leonow32/verilog-fpga>

- **1 bias** – stosowane są tylko dwa stany: 0 oraz 1,
- **1/2 bias** – stosowane są trzy stany: 0, 1/2 i 1,
- **1/3 bias** – stosowane są cztery stany: 0, 1/3, 2/3 i 1,
- **1/4 bias** – stosowanych jest pięć stanów: 0, 1/4, 1/2, 3/4 oraz 1,
- i tak dalej...

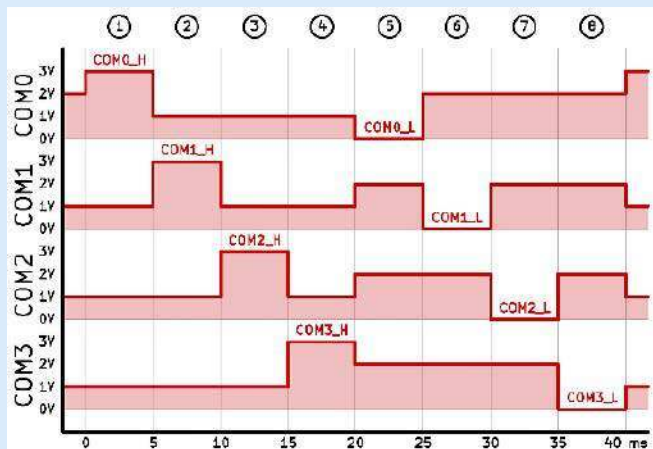
Wyświetlacz, jaki będziemy stosować, ma **1/4 duty** oraz **1/3 bias**, zatem przy napięciu zasilania 3,3 V poszczególne stany będą mieć napięcia 0 V; 1,1 V; 2,2 V oraz 3,3 V. Dla wygody w dalszej części tekstu część ułamkową będziemy zaokrąglać.

Aby było jeszcze trudniej – między każdym segmentem a elektrodami wspólnymi **napięcie średnie** musi być równe zero. Jeżeli doprowadzimy do tego, że segmenty wyświetlacza będą zasilane napięciem stałym, to po pewnym czasie dojdzie do uszkodzenia ciekłych kryształów. Zatem musimy w bardzo specyficzny sposób naprzemiennie odwracać sygnały sterujące wyświetlaczem w taki sposób, żeby **napięcie chwilowe** pomiędzy elektrodami COM i SEG było równe +3 V lub –3 V dla segmentów widocznych oraz +1 V lub –1 V dla segmentów niewidocznych.

Zobaczmy teraz, jak wyglądają przebiegi napięć na elektrodach wspólnych COM naszego wyświetlacza, które pokazano na **rysunku 2**. Takie przebiegi muszą występować przez cały czas pracy niezależnie od tego, które segmenty wyświetlacza mają być zaczernione.



Rysunek 1. Schemat połączeń wewnątrz wyświetlacza LCD



Rysunek 2. Przebiegi na elektrodach COM

Cały cykl sterowania można podzielić na osiem części, które łącznie stanowią jedną ramkę. Każda z tych części trwa 5 milisekund. Części możemy podzielić na dwie grupy:

- W częściach od 1 do 4 – aktywna elektroda COM ma napięcie 3 V, a nieaktywne elektrody COM mają napięcia 1 V.
- W częściach od 5 do 8 – aktywna elektroda COM ma napięcie 0 V, a nieaktywne elektrody COM mają napięcie 2 V.

W ten sposób aktywna jest zawsze tylko jedna z czterech elektrod COM. Natychmiast po zmianie aktywnej elektrody COM musimy aktywować odpowiednie elektrody SEG, aby zaczernić żądane segmenty.

Na rysunku 3 pokazano przykład sygnałów, które powodują zaczernienie segmentów C i D cyfry zerowej. Segment C zaczerniany jest, kiedy aktywna jest para COM1-SEG0, a segment D staje się widoczny po aktywacji pary COM0-SEG1. Dla poprawy czytelności pominięto nieistotne sygnały COM2, COM3 oraz SEG2...7.

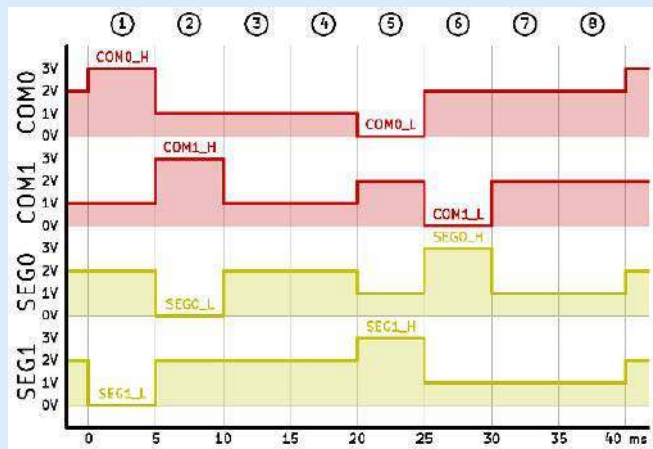
Podobnie jak w przypadku sterowaniu elektrodami COM, sterowanie elektrodami SEG możemy podzielić na dwie grupy:

- W częściach od 1 do 4 – aktywna elektroda SEG ma napięcie 0 V, a nieaktywne elektrody SEG mają napięcia 2 V.
- W częściach od 5 do 8 – aktywna elektroda SEG ma napięcie 3 V, a nieaktywne elektrody SEG mają napięcie 1 V.

W tabeli 1 zestawiono powyższe informacje w sposób bardziej syntetyczny dla wszystkich możliwych kombinacji aktywnych/nieaktywnych elektrod COM/SEG. Tylko w przypadku, kiedy aktywne są jednocześnie COM i SEG, na wybranym segmencie występuje napięcie ± 3 V, co powoduje zaczernienie segmentu, a we wszystkich innych przypadkach napięcie segmentu wynosi ± 1 V i w rezultacie segment pozostaje niewidoczny.

Generowanie napięć sterujących

Zastanówmy się teraz, skąd wziąć napięcia o wartości 0, 1/3, 2/3 i 1 napięcia zasilającego? Opcje są dwie. Schemat pierwszego rozwiązania został pokazany na rysunku 4. Jest to rozwiązanie analogowe i polega na zastosowaniu dzielnika napięcia z trzech rezystorów o identycznej rezystancji, które tworzą napięcia o wartości 1/3 i 2/3 napięcia



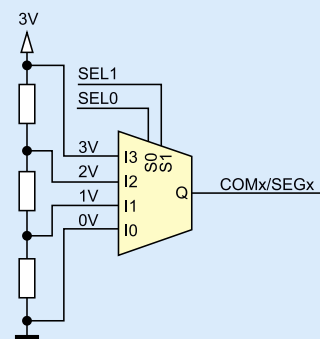
Rysunek 3. Przykład sygnałów powodujących zaczernienie segmentów C i D cyfry zerowej

zasilającego. Stan 0 to oczywiście podłączenie do masy, a 1 to połączenie z zasilaniem. Następnie za pomocą multiplexera analogowego wybierane jest jedno z czterech dostępnych napięć. Takich multiplexerów potrzebujemy 12, czyli tyle, ile pinów ma wyświetlacz. Wyjścia multiplexerów wychodzą na piny układu scalonego, a one są połączone bezpośrednio z elektrodami wyświetlacza LCD.

Niestety w FPGA nie mamy do dyspozycji dzielników napięcia ani multiplexerów analogowych. Generowanie różnych napięć musimy zrealizować całkowicie cyfrowo. Rozwiązaniem jest generator PWM oraz filtry RC. Rozwiązanie pokazano na rysunku 5. Dwanaście sygnałów PWM, pochodzących z pinów FPGA, przechodzi przez dwanaście filtrów RC, które przekształcają je na sygnały analogowe.

Generatorom PWM przyjrzymy się bliżej w jednym z kolejnych odcinków kursu. Jedyne, co musimy o nich wiedzieć, to tylko to, że w naszym rozwiązaniu będą one generować sygnał prostokątny o wypełnieniu 33% lub 66%, co po przefiltrowaniu da napięcie około 1 V i 2 V.

Filtry RC można pominąć. Wtedy sygnały PWM zostaną wygładzone przez pasywną pojemność segmentów wyświetlacza. Jednak takie rozwiązanie nie jest rekomendowane ze względu na generowanie dużo większych zakłóceń, przez co urządzenie może mieć



Rysunek 4. Fragment analogowego sterownika pinów wyświetlacza LCD

Tabela 1. Zestawienie napięć na aktywnych i nieaktywnych elektrodach COM i SEG

Części ramki	Aktywny COM	Aktywny SEG	Różnica	Części ramki	Aktywny COM	Nieaktywny SEG	Różnica
1 2 3 4	3 V	0 V	+3 V	1 2 3 4	3 V	2 V	+1 V
5 6 7 8	0 V	3 V	-3 V	5 6 7 8	0 V	1 V	-1 V
Części ramki	Nieaktywny COM	Aktywny SEG	Różnica	Części ramki	Nieaktywny COM	Nieaktywny SEG	Różnica
1 2 3 4	1 V	0 V	+1 V	1 2 3 4	1 V	2 V	-1 V
5 6 7 8	2 V	3 V	-1 V	5 6 7 8	2 V	1 V	+1 V

problemy z certyfikacją EMC. Zaleca się umieścić filtry w pobliżu pinów układu FPGA, aby ścieżki sygnałów PWM były jak najkrótsze.

Na **listingu 1** zaprezentowano kod modułu odpowiedzialnego za generowanie czterech napięć, jakie są wymagane do sterowania wyświetlaczem LCD, którego będziemy używać. Postanowiłem zmodyfikować nieco konwencję pisania kodu, jaką stosowałem w poprzednich odcinkach ze względu na to, że nasze projekty stają się coraz bardziej skomplikowane.

Pierwsza zmiana polega na zastosowaniu instrukcji ``default_nettype` w linii 1. W języku Verilog nie ma potrzeby deklarowania zmiennych przed ich użyciem, tak jak to jest w C++ czy wielu innych językach. Może się zdarzyć, że popełniając literówkę w nazwie zmiennej, syntezyzator stwierdzi, że chcemy utworzyć zupełnie nową zmienną typu `wire`. W rezultacie kod będzie poprawny pod względem składni języka i zsyntezuje się, ale nie będzie działał prawidłowo. Rozwiązaniem tego problemu jest zastosowanie rozwiązania z linii 1. W momencie napotkania zmiennej, która nie została wcześniej zdefiniowana jako zmienna `wire`, `reg`, `integer` czy innego typu, zostaje zgłoszony błąd.

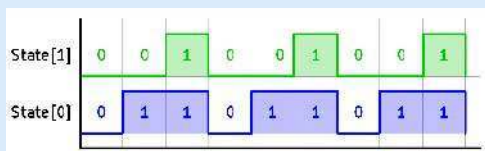
Instrukcja ta ma zasięg globalny i dotyczy wszystkich plików projektu, więc mogłaby spowodować błąd w plikach, które napisane są „po staremu”. Z tego powodu w ostatniej linijce pliku przywracamy ustawienie, że domyślnie wybieranym typem dla niezdefiniowanych wcześniej zmiennych jest typ `wire` (linia 10).

Kolejna zmiana polega na dodawaniu oznaczenia do nazwy zmiennej w zależności od kierunku portu lub typu. Do wszystkich wyjść będziemy dodawać `_o` jak `output`, a do wejść będzie to `_i` jak `input`. W ten sposób będziemy mogli rozróżnić wejścia i wyjścia, analizując instancję w module nadrzędnym. Wyjątkiem są `Clock` i `Reset`, ponieważ one występują w prawie każdym module i ich przeznaczenie jest oczywiste. Ponadto do nazw zmiennych typu `reg` dodamy `_r`, a do zmiennych `wire` dodawać będziemy `_w`.

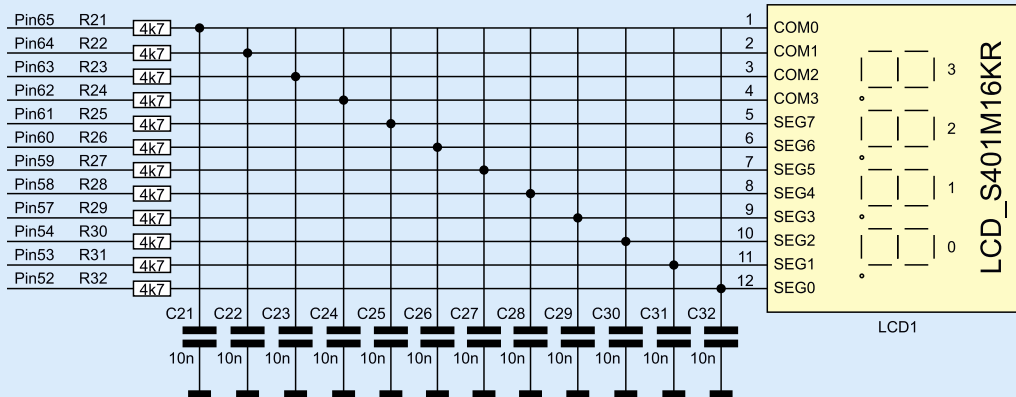
W jaki sposób będziemy generować przebiegi o wypełnieniu 33% i 66%? Potrzebujemy dwubitowej zmiennej `State_r` (linia 2). W każdym taktie zegara będziemy tę zmienną modyfikować, by uzyskać takie przebiegi, jak pokazano na **rysunku 6**.

W bloku jedynym `always` tego modułu mamy proste drzewko decyzyjne, które na podstawie obecnego stanu zmiennej `State_r` ustawia kolejny stan. I tak, jeżeli obecna wartość tej zmiennej jest równa 00, to zmieniamy ją na 01 (linia 3). Jeżeli już mamy 01, to zmieniamy na 11 (linia 4), a w innym przypadku zmieniamy z powrotem na 00. W ten sposób uzyskujemy cykliczną, 3-cykłową pracę, gdzie `State[1]` przez dwa takty ma stan niski i przez jeden takt ma stan wysoki. Natomiast `State[0]` przez dwa takty jest w stanie wysokim, a tylko przez jeden – w niskim.

Pozostaje już tylko przypisać odpowiednie sygnały do portów wyjściowych. W linii 6 przypisujemy wyjście, które ma dostarczać napięcie 0 V, czyli przywiązujemy je na stałe ze stanem niskim.



Rysunek 6. Przebiegi PWM o wypełnieniu 33% i 66%



Rysunek 5. Schemat połączenia wyświetlacza LCD z układem FPGA poprzez filtry RC

Listing 1. Kod pliku `lcd_pwm.v`

```
// Plik lcd_pwm.v

`default_nettype none // 1
module LCD_PWM(
    input wire Clock,
    input wire Reset,
    output wire Voltage0_o, // Wypełnienie 0%
    output wire Voltage1_o, // Wypełnienie 33%
    output wire Voltage2_o, // Wypełnienie 66%
    output wire Voltage3_o // Wypełnienie 100%
);

// Bardzo prosta maszyna stanów // 2
reg [1:0] State_r;
always @(posedge Clock, negedge Reset) begin
    if (!Reset)
        State_r <= 2'b00;
    else if (State_r == 2'b00)
        State_r <= 2'b01; // 3
    else if (State_r == 2'b01)
        State_r <= 2'b11; // 4
    else
        State_r <= 2'b00; // 5
end

// Przypisanie wyjść // 6
assign Voltage0_o = 1'b0; // 7
assign Voltage1_o = State_r[1]; // 8
assign Voltage2_o = State_r[0]; // 9
assign Voltage3_o = 1'b1; // 10

endmodule
`default_nettype wire // 10
```

Analogicznie postępujemy z wyjściem, które ma dostarczać 3 V – przypisujemy do niego stan wysoki (linia 9). Do wyjścia, mającego dostarczać napięcie 1 V, przywiązujemy rejestr `State_r[1]` (linia 7), a do wyjścia 2 V dajemy `State_r[0]` (linia 8).

Warto zwrócić uwagę, że sygnały generowane przez moduł z listingu 1 zmieniają się z każdym taktom zegara, który może mieć częstotliwość powyżej 100 MHz. Takie szybkie przełączanie pinów GPIO nie jest optymalne – może powodować dużo większe zużycie energii, niż jest to potrzebne, a także może utrudniać zaliczenie testów EMC. Do tego tematu powrócimy w odcinku poświęconym generatorom PWM, a na razie pozostawmy tę małą niedoskonałość.

Sterownik wyświetlacza LCD

Przejdźmy teraz do modułu, który decyduje o tym, jakie napięcie ma zostać ustalone na poszczególnych elektrodach wyświetlacza na podstawie informacji o tym, jakie segmenty mają być widoczne. Moduł ten widać na **listingu 2**. Standardowo, zaczynamy od parametru, który określa częstotliwość sygnału zegarowego `CLOCK_HZ` (linia 1). Drugi parametr nazwany `CHANGE_COM_US` określa czas w mikrosekundach, po upływie którego przełączane są elektrody wspólne COM (linia 2). W naszym przykładzie ten czas to 5 ms. Ponieważ w cyklu sterowania elektrodami wspólnymi jest 8 stanów, gdzie każdy trwa 5 ms, całość zajmie 40 ms. Odwrotność tej liczby daje nam 25 Hz. Jest to częstotliwość odświeżania wyświetlacza.

W linii 3 i kolejnych tworzymy cztery 8-bitowe wejścia, sterujące segmentami wszystkich cyfr. Cyfry ponumerowane są zgodnie z opisem na rysunku 1. Najmłodszy bit każdego z wejść `DigitX_i[0]`

Listing 2. Kod pliku lcd.v

```

// Plik lcd.v

`default_nettype none
module LCD #(
    parameter CLOCK_HZ      = 10_000_000,           // 1
    parameter CHANGE_COM_US = 5000                 // 2
)(
    input wire Clock,
    input wire Reset,
    input wire [7:0] Digit3_i, // PGFEDCBA           // 3
    input wire [7:0] Digit2_i,
    input wire [7:0] Digit1_i,
    input wire [7:0] Digit0_i,
    output wire [3:0] ComPWM_o,                       // 4
    output wire [7:0] SegPWM_o                       // 5
);

// Generator czterech napięć zasilających elektrody wyświetlacza
wire [3:0] Voltage_w; // 6
LCD_PWM LCD_PWM_inst(
    .Clock(Clock), // 7
    .Reset(Reset),
    .Voltage0_o(Voltage_w[0]),
    .Voltage1_o(Voltage_w[1]),
    .Voltage2_o(Voltage_w[2]),
    .Voltage3_o(Voltage_w[3])
);

// Generator sygnałów przełączających stan co określony czas
wire ChangeState_w; // 8
StrobeGenerator #(
    .CLOCK_HZ(CLOCK_HZ), // 9
    .PERIOD_US(CHANGE_COM_US)
) StrobeGenerator0(
    .Clock(Clock),
    .Reset(Reset),
    .Strobe_o(ChangeState_w)
);

// Maszyna stanów // 10
reg [2:0] State_r /* synthesis syn_encoding = "safe, sequential"
*/;
localparam COM_0H = 3'd0;
localparam COM_1H = 3'd1;
localparam COM_2H = 3'd2;
localparam COM_3H = 3'd3;
localparam COM_0L = 3'd4;
localparam COM_1L = 3'd5;
localparam COM_2L = 3'd6;
localparam COM_3L = 3'd7;

// Zmiana stanu
always @(posedge Clock, negedge Reset) begin
    if(!Reset)
        State_r <= 0;
    else if(ChangeState_w) // 11
        State_r <= State_r + 1'b1; // 12
end

// Zmienne przechowujące aktualne napięcie elektrod wyświetlacza
reg [1:0] ComAnalog_r[0:3]; // 13
reg [1:0] SegAnalog_r[0:7]; // 14

// Ustalanie napięcia elektrod w zależności od stanu maszyny
// oraz informacji o tym, które segmenty mają być widoczne
always @(*) begin // 15
    case(State_r) // 16
        COM_0H: begin
            ComAnalog_r[0] = 2'd3;
            ComAnalog_r[1] = 2'd1;
            ComAnalog_r[2] = 2'd1;
            ComAnalog_r[3] = 2'd1;
            SegAnalog_r[0] = Digit0_i[7] ? 2'd0 : 2'd2; //

Dwukropek
            SegAnalog_r[1] = Digit0_i[3] ? 2'd0 : 2'd2; // D0, D
            SegAnalog_r[2] = Digit1_i[7] ? 2'd0 : 2'd2; // D1, P
            SegAnalog_r[3] = Digit1_i[3] ? 2'd0 : 2'd2; // D1, D
            SegAnalog_r[4] = Digit2_i[7] ? 2'd0 : 2'd2; // D2, P
            SegAnalog_r[5] = Digit2_i[3] ? 2'd0 : 2'd2; // D2, D
            SegAnalog_r[6] = Digit3_i[7] ? 2'd0 : 2'd2; // D3, P
            SegAnalog_r[7] = Digit3_i[3] ? 2'd0 : 2'd2; // D3, D
        end

        COM_1H: begin
            ComAnalog_r[0] = 2'd1;
            ComAnalog_r[1] = 2'd3;
            ComAnalog_r[2] = 2'd1;
            ComAnalog_r[3] = 2'd1;
            SegAnalog_r[0] = Digit0_i[2] ? 2'd0 : 2'd2; // D0, C
            SegAnalog_r[1] = Digit0_i[4] ? 2'd0 : 2'd2; // D0, E
            SegAnalog_r[2] = Digit1_i[2] ? 2'd0 : 2'd2; // D1, C
            SegAnalog_r[3] = Digit1_i[4] ? 2'd0 : 2'd2; // D1, E
            SegAnalog_r[4] = Digit2_i[2] ? 2'd0 : 2'd2; // D2, C
            SegAnalog_r[5] = Digit2_i[4] ? 2'd0 : 2'd2; // D2, E
            SegAnalog_r[6] = Digit3_i[2] ? 2'd0 : 2'd2; // D3, C
            SegAnalog_r[7] = Digit3_i[4] ? 2'd0 : 2'd2; // D3, E
        end

        COM_2H: begin
            ComAnalog_r[0] = 2'd1;
            ComAnalog_r[1] = 2'd1;
            ComAnalog_r[2] = 2'd3;
            ComAnalog_r[3] = 2'd1;
            SegAnalog_r[0] = Digit0_i[1] ? 2'd0 : 2'd2; // D0, B
            SegAnalog_r[1] = Digit0_i[5] ? 2'd0 : 2'd2; // D0, F
            SegAnalog_r[2] = Digit1_i[1] ? 2'd0 : 2'd2; // D1, B
            SegAnalog_r[3] = Digit1_i[5] ? 2'd0 : 2'd2; // D1, F
            SegAnalog_r[4] = Digit2_i[1] ? 2'd0 : 2'd2; // D2, B
            SegAnalog_r[5] = Digit2_i[5] ? 2'd0 : 2'd2; // D2, F
            SegAnalog_r[6] = Digit3_i[1] ? 2'd0 : 2'd2; // D3, B
            SegAnalog_r[7] = Digit3_i[5] ? 2'd0 : 2'd2; // D3, F
        end

        COM_3H: begin
            ComAnalog_r[0] = 2'd1;
            ComAnalog_r[1] = 2'd1;
            ComAnalog_r[2] = 2'd1;
            ComAnalog_r[3] = 2'd1;
            SegAnalog_r[0] = Digit0_i[0] ? 2'd0 : 2'd2; // D0, A
            SegAnalog_r[1] = Digit0_i[6] ? 2'd0 : 2'd2; // D0, G
            SegAnalog_r[2] = Digit1_i[0] ? 2'd0 : 2'd2; // D1, A
            SegAnalog_r[3] = Digit1_i[6] ? 2'd0 : 2'd2; // D1, G
            SegAnalog_r[4] = Digit2_i[0] ? 2'd0 : 2'd2; // D2, A
            SegAnalog_r[5] = Digit2_i[6] ? 2'd0 : 2'd2; // D2, G
            SegAnalog_r[6] = Digit3_i[0] ? 2'd0 : 2'd2; // D3, A
            SegAnalog_r[7] = Digit3_i[6] ? 2'd0 : 2'd2; // D3, G
        end
    endcase
end

// Przypisanie wyjść // 18
assign ComPWM_o[0] = Voltage_w[ComAnalog_r[0]];
assign ComPWM_o[1] = Voltage_w[ComAnalog_r[1]];
assign ComPWM_o[2] = Voltage_w[ComAnalog_r[2]];
assign ComPWM_o[3] = Voltage_w[ComAnalog_r[3]];
assign SegPWM_o[0] = Voltage_w[SegAnalog_r[0]];
assign SegPWM_o[1] = Voltage_w[SegAnalog_r[1]];
assign SegPWM_o[2] = Voltage_w[SegAnalog_r[2]];
assign SegPWM_o[3] = Voltage_w[SegAnalog_r[3]];
assign SegPWM_o[4] = Voltage_w[SegAnalog_r[4]];
assign SegPWM_o[5] = Voltage_w[SegAnalog_r[5]];
assign SegPWM_o[6] = Voltage_w[SegAnalog_r[6]];
assign SegPWM_o[7] = Voltage_w[SegAnalog_r[7]];

endmodule
`default_nettype wire

```

odpowiada za sterowanie segmentem A wyświetlacza, natomiast bit najstarszy **DigitX_i[7]** odpowiada za segment P, czyli przecinek lub dwukropek. Stan wysoki powodować będzie zaczernienie odpowiadającego segmentu, a stan niski sprawi, że segment będzie niewidoczny.

Linia 4 zawiera wyjście sygnałów PWM sterujących czterema elektrodami wspólnymi COM, a w linii 5 jest 8-bitowe wyjście kontrolujące elektrody SEG.

Następnie tworzymy instancję modułu, który omawialiśmy w poprzednich akapitach (linia 7). Sygnały PWM o wypełnieniu 0%, 33%, 66% i 100% są rozprowadzane do pozostałych elementów za pomocą 4-bitowej zmiennej **Voltage_w** typu wire (linia 6). Zgrupowanie sygnałów PWM w zmienną 4-bitową będzie bardzo pomocne w dalszej części kodu – stosując wyrażenie **Voltage_w[X]**, gdzie X to liczby od 0 do 3, będziemy mogli w bardzo łatwy sposób uzyskać sygnał generujący napięcie od 0 do 3 woltów. Będziemy wykorzystywać zależność, że indeks w nawiasach kwadratowych jest równy napięciu w woltach. Będzie to omówione w dalszej części tekstu.

W dalszej części kodu tworzymy instancję generatora sygnałów strobe, które służyć mają do przełączania maszyny stanów. W linii 9 stosujemy moduł **StrobeGenerator**, który znamy z poprzednich odcinków kursu i nie będziemy po raz kolejny omawiać jego działania. Sygnał zmiany stanu jest przekazywany za pośrednictwem zmiennej wire **ChangeState_w** (linia 8) – sygnał ten pozostaje w stanie niskim przez większość czasu, ale cyklicznie co 5 ms jest ustawiany w stan wysoki na jeden takt zegara, co ma spowodować przełączenie zmiennej sterującej maszyną stanów.

Maszynę stanów tworzymy w linii 10 i kolejnych. Zgodnie z rysunkiem 2 i 3 wyświetlacz ma osiem stanów – z tego powodu zmienna **State_r**, przechowująca stan maszyny, jest zmienną 3-bitową. W następnych liniach nazywamy wszystkie stany poprzez tworzenie parametrów lokalnych, którym przypisane są 3-bitowe liczby od 0 do 7.

Zapis `/* synthesis syn_encoding = "safe, sequential" */` jest czymś zupełnie nowym. Z punktu widzenia składni języka Verilog jest to zupełnie zwyczajny komentarz, jednak dla syntezy Lattice Synthesis Engine jest to polecenie mówiące, w jaki sposób ma zostać zrealizowana maszyna stanów. Parametr **syn_encoding** może przyjmować następujące wartości:

- **sequential** – stan maszyny zapisany jest binarnie, a stany ponumerowane są od zera; zmienna sterująca maszyną potrzebuje tyle bitów, ile potrzebne jest, by “zmieścić” numer ostatniego stanu. Takie kodowanie oszczędza zasoby, zwłaszcza przerzutniki, ale na ogół działa wolniej niż one-hot;
- **one-hot** – stan maszyny zapisany jest za pomocą kodowania one-hot, tzn. zmienna przechowująca stan ma tyle bitów, ile możliwych jest stanów i tylko jeden z nich jest równy 1, a pozostałe bity są równe 0. Taki sposób przyspiesza działanie, lecz jest bardziej zasobochłonny;
- **gray** – rozwiązanie pośrednie, można stosować, tylko gdy maszyna ma nie więcej niż 4 stany;
- **safe** – może być dodane opcjonalnie do sequential i one-hot; zostanie dodatkowo dodana logika resetująca maszynę stanów w przypadku, gdyby zmienna stanu miała wartość, która nie jest obsługiwana (np. dwa bity równe 1 przy kodowaniu one-hot).

Zachęcam, by przetestować kod dla różnych ustawień maszyny stanów. Lattice Synthesis Engine automatycznie przerobi nasz kod i nie ma potrzeby ręcznie zmieniać definicji stanów. W tabeli 2 porównano wyniki, jakie są osiągnięte dla różnych ustawień.

W liniach 13 i 14 mamy coś nowego, czego jeszcze w tym kursie nie widzieliśmy – tablice, zwane także pamięciami. Pierwszy nawias kwadratowy informuje, ile bitów ma pojedyncza zmienna tablicy, dokładnie w taki sam sposób, jak przy wielobitowych zmiennych reg i wire, które już znamy. Zatem w liniach 13 i 14 tworzymy tablice zmiennych 2-bitowych. Drugi nawias kwadratowy informuje

Tabela 2. Porównanie parametrów różnych implementacji maszyny stanów

Implementacja	Max. częstotliwość	Przerzutniki	Slice	LUT4
One-hot (domyślny)	116 MHz	66	88	176
One-hot safe	105 MHz	68	94	187
Sequential	123 MHz	61	69	132

o numerach indeksu pierwszego i ostatniego elementu tablicy. Tablica **ComAnalog_r** ma cztery elementy 2-bitowe, ponumerowane od 0 do 3, a tablica **SegAnalog_r** ma osiem elementów 2-bitowych, ponumerowanych od 0 do 7.

Jak zapewne się spodziewasz, w tych tablicach będziemy przechowywać informację o napięciu, jakie ma zostać dostarczone do odpowiednich elektrod COM i SEG. Zgrupowanie ich w tablice nie jest konieczne – równie dobrze moglibyśmy utworzyć dwanaście zwykłych 2-bitowych zmiennych typu reg, ale taki zabieg poprawia czytelność kodu i umożliwia jego parametryzację. W sytuacji, gdybyśmy pisali moduł obsługujący wyświetlacze o różnej liczbie elektrod COM/SEG, zastosowanie tablic byłoby zdecydowanym ułatwieniem.

Blok maszyny stanów znajdziemy w linii 15. Jest to blok logiki kombinacyjnej, reagujący na zmianę dowolnego sygnału występującego w bloku – poznajemy to po znaku gwiazdki za **always @(*)**. Składa się on z dość rozbudowanej instrukcji case, która decyduje, co ma się wykonywać na podstawie aktualnego stanu zmiennej **State_r**. Dalszy kod podzielony jest na osiem części dla każdego ze stanów, jakie utworzyliśmy w linii 10.

Przeanalizujemy tylko stan **COM_0H** – wszystkie są zrealizowane bardzo podobnie i wynikają z konieczności wygenerowania przebiegów, jakie pokazano na rysunkach 2 i 3. Jest to stan, w którym elektroda COM0 ma mieć napięcie 3 V, a pozostałe COM-y mają mieć 1 V. Z tego powodu do **ComAnalog_r[0]** wpisujemy 2-bitową wartość 2'd3, a do pozostałych 2'd1;

Z ustaleniem napięć na elektrodach SEG sprawa jest bardziej skomplikowana, ponieważ zależy od poszczególnych bitów wejść **DigitX_i**. Dla każdego z pinów **SegAnalog_r[X]** musimy sprawdzić właściwy bit na odpowiednim wejściu danych (1 – segment widoczny, 0 – segment niewidoczny) i najlepiej jest to zrobić za pomocą operatora warunkowego ?: który występuje również w językach C i C++. Jeżeli sprawdzany bit ma stan wysoki, to do zmiennej **SegAnalog_r[X]** wpisujemy 2'd0, a jeżeli niski, to 2'd2. Analogicznie postępujemy we wszystkich pozostałych siedmiu stanach.

Przeskoczmy teraz do linii 18 i kolejnych. Przypisujemy tutaj wyjścia, które sterują pinami wyświetlacza. Poszczególne wyjścia **CompPWM_o** oraz **SegPWM_o** podłączamy do jednego z czterech sygnałów **Voltage_w**. Robimy to za pomocą operatora **[]**. Taka konstrukcja zostanie zsyntezowana jako multiplexer z czterema wejściami danych i 2-bitowym wejściem adresowym, tzn. adres wybierający może być równy 0, 1, 2 lub 3.

Te liczby bezpośrednio odpowiadają napięciu, jakie zostanie wygenerowane po przefiltrowaniu sygnału PWM. Pobieramy je z tablic **ComAnalog_r** i **SegAnalog_r**, które zawierają 2-bitowe zmienne typu reg. W ten prosty sposób zmiana jednej zmiennej w tablicy spowoduje zmianę napięcia na wybranej elektrodzie wyświetlacza.

Zwróć uwagę, że indeksy w nawiasach kwadratowych we wszystkich przypisaniach **SegPWM_o[x]** oraz **SegAnalog_r[x]** są sobie równe. Gdyby zaszła potrzeba napisania sterownika, który obsługuje dowolną liczbę elektrod, można by wtedy ten kod łatwo sparаметryzować. To samo dotyczy **CompPWM[x]** oraz **ComAnalog_r[x]**.

Testbench sterownika

Czas, aby zrobić testbench, którym będziemy mogli symulować sterownik wyświetlacza LCD. Jego kod został pokazany na listingu 3.

Aby zredukować czas trwania symulacji oraz ilość danych w pliku wynikowym, ustawimy częstotliwość symulowanego zegara na 1 MHz (linia 1), a czas trwania jednego stanu wyświetlacza ustawimy na 50 mikrosekund (linia 5).

W linii 2 mamy cztery 8-bitowe zmienne, które symulują dane wejściowe dla sterownika wyświetlacza. Każdy z bitów tych zmiennych odpowiedzialny jest za inny segment wyświetlacza. Celowo jeden z nich jest w stanie wysokim – dzięki temu będziemy mogli porównać przebiegi sygnałów dla segmentów widocznych i niewidocznych.

Napięcia na elektrodach wyświetlacza i ich zmiany w czasie wyświetlimy w formie tabelarycznej. W tym celu w linii 3 wyświetlamy nagłówek tabeli za pomocą funkcji **\$display**. Następnie mamy funkcję **\$monitor**, która powoduje wyświetlenie tekstu, kiedy zmienia się jakakolwiek z obserwowanych zmiennych. Obie funkcje działają podobnie do *printf()* z C++. W swoim pierwszym argumencie przyjmuje ciąg znaków, który ma być wyświetlany, w miejsce **%d** mają być wstawione zmienne liczbowe, a w miejscu **%t** ma być wyświetlony czas. W kolejnych argumentach znajdują się zmienne, odpowiadające znakom **%** w ciągu.

W linii 4 umieszczono pętlę **repeat**, która wykona się osiem razy. Ta liczba jest nie bez powodu. Właśnie tyle stanów jest w całym cyklu pracy wyświetlacza. Wewnątrz pętli znajduje się tylko jedno polecenie – oczekiwanie na wystąpienie zbocza rosnącego sygnału **ChangeState_w** wewnątrz instancji **DUT**, czyli testowanego sterownika wyświetlacza. Stan wysoki na tym sygnale oznacza żądanie przejścia do kolejnego stanu. W taki sposób zapewniamy, że symulacja będzie trwała przez osiem stanów – niezależnie od częstotliwości sygnału zegarowego ani czasu trwania stanu, który jest ustalany parametrem **CHANGE_COM_US**.

Następnie mamy instancję sterownika wyświetlacza. Nie ma w niej nic nowego, co by wymagało komentarza.

Dalej znajduje się blok *initial*, w którym są tylko instrukcje związane z eksportem danych, powstałych w wyniku symulacji. Instrukcja **\$dumpvars()**, której pierwszy argument jest zerem, powoduje zapisywanie wszystkich zmiennych z modułu wskazanego w drugim argumencie i wszystkich zmiennych w jego modułach podrzędnych. Ale to nie dotyczy tablic! Niestety musimy zażądać eksportu każdej zmiennej tablicy osobno, tak jak to przedstawiono w linii 6. W przypadku większych tablic można posłużyć się pętlą **for**.

Symulacja

Kod skryptu, uruchamiającego symulację w symulatorze Icarus Verilog, pokazano na **listingu 4**. Wystarczy dodać ten skrypt do Lattice Diamond i uruchomić go.

W wyniku symulacji dostajemy plik *lcd.vcd*. Należy go otworzyć. Uruchomi się przeglądarka GTKWave. Omawialiśmy ją dokładnie w 12 odcinku kursu. Dodaj i skonfiguruj wszystkie sygnały tak, aby uzyskać efekt pokazany na **rysunku 7**. W razie potrzeby gotowy plik GTWK z kompletną analizą jest dostępny w plikach z kodem źródłowym.

Zmienna **State_r** prezentuje stan sterownika wyświetlacza. Dla poprawy czytelności wykresu warto zmienić wartości liczbowe od 0 do 7 na etykiety tekstowe, które jasno i czytelnie będą nam mówić, w jakim stanie znajduje się maszyna stanów. W tym celu musimy utworzyć dodatkowy plik z etykietami. Nazwijmy go *lcd-state.gtkw*. Treść tego pliku pokazano na **listingu 5**.

Aby wyświetlić etykiety stanów, zmienną **State_r** w okienku **Signals** klikamy prawym przyciskiem myszy, wybieramy opcję **Data Format**, następnie **Translate Filter File** i ostatecznie **Enable and Select**. Pojawi się okienko, w którym klikamy przycisk **Add**

Listing 3. Kod pliku *lcd_tb.v*

```
// Plik lcd_tb.v

`timescale 1ns/1ns
`default_nettype none
module LCD_tb();

    parameter CLOCK_HZ      = 1_000_000;           // 1
    parameter HALF_PERIOD_NS = 1_000_000_000 / (2 * CLOCK_HZ);

    // Generator sygnału zegarowego
    reg Clock = 1'b1;
    always begin
        #HALF_PERIOD_NS;
        Clock = !Clock;
    end

    // Zmienne
    reg      Reset      = 1'b0;
    reg [7:0] Digit3_r = 8'b00000000;           // 2
    reg [7:0] Digit2_r = 8'b00000000;
    reg [7:0] Digit1_r = 8'b00010000; // widoczny segment E
    reg [7:0] Digit0_r = 8'b00000000;
    wire [3:0] ComPWM_w;
    wire [7:0] SegPWM_w;

    // Sekwencja testowa
    initial begin
        $timeformat(-6, 3, "us", 10);
        $display("===== START =====");
        $display("CLOCK_HZ = %9d", CLOCK_HZ);

        #1 Reset = 1'b1;

        $display("time C0 C1 C2 C3 S0 S1 S2 S3 S4 S5 S6 S7"); // 3
        $monitor("%t %d %d %d %d %d %d %d %d %d %d %d %d",
            $realtime,
            DUT.ComAnalog_r[0],
            DUT.ComAnalog_r[1],
            DUT.ComAnalog_r[2],
            DUT.ComAnalog_r[3],
            DUT.SegAnalog_r[0],
            DUT.SegAnalog_r[1],
            DUT.SegAnalog_r[2],
            DUT.SegAnalog_r[3],
            DUT.SegAnalog_r[4],
            DUT.SegAnalog_r[5],
            DUT.SegAnalog_r[6],
            DUT.SegAnalog_r[7],
        );

        repeat(8) begin                               // 4
            @(posedge DUT.ChangeState_w);
        end

        #1 $display("===== END =====");
        #1 $finish;
    end

    // Instancja testowanego modułu
    LCD #(
        .CLOCK_HZ(CLOCK_HZ),
        .CHANGE_COM_US(50)                             // 5
    ) DUT(
        .Clock(Clock),
        .Reset(Reset),
        .Digit3_i(Digit3_r),
        .Digit2_i(Digit2_r),
        .Digit1_i(Digit1_r),
        .Digit0_i(Digit0_r),
        .CompPWM_o(ComPWM_w),
        .SegPWM_o(SegPWM_w)
    );

    // Eksport zmiennych
    initial begin
        $dumpfile("lcd.vcd");
        $dumpvars(0, LCD_tb);
        $dumpvars(2, DUT.ComAnalog_r[0]);           // 6
        $dumpvars(2, DUT.ComAnalog_r[1]);
        $dumpvars(2, DUT.ComAnalog_r[2]);
        $dumpvars(2, DUT.ComAnalog_r[3]);
        $dumpvars(2, DUT.SegAnalog_r[0]);
        $dumpvars(2, DUT.SegAnalog_r[1]);
        $dumpvars(2, DUT.SegAnalog_r[2]);
        $dumpvars(2, DUT.SegAnalog_r[3]);
        $dumpvars(2, DUT.SegAnalog_r[4]);
        $dumpvars(2, DUT.SegAnalog_r[5]);
        $dumpvars(2, DUT.SegAnalog_r[6]);
        $dumpvars(2, DUT.SegAnalog_r[7]);
    end

endmodule
`default_nettype wire
```

filter to list i wskazujemy plik *lcd-state.gtkw*. Ścieżka do pliku pojawi się w okienku. Należy ją kliknąć, aby była podświetlona (inaczej nie będzie działać!), po czym możemy kliknąć **OK**. Nic się nie zmieniło! Musimy jeszcze raz kliknąć prawym przyciskiem myszy zmienną **State_r** w okienku **Signals**, po czym wybieramy jeszcze raz **Data Format** i tym razem klikamy **Enum** (autorowi programu należą się szczerze gratulacje za prosty i intuicyjny interfejs).

Zwróć uwagę na zmienną **SegPWM_o[3]**. Jest to stan na wyprowadzeniu FPGA, który steruje elektrodą SEG3 wyświetlacza poprzez filtr RC. W stanie COM1_H oraz COM1_L ten sygnał wygląda zupełnie inaczej niż wszystkie inne. To jest właśnie sygnał, który powoduje zaczernienie segmentu E cyfry 1. Wynika on z ustawienia piątego bitu zmiennej **Digit1_r** (listing 3, linia 2). Zgodnie z rysunkiem 1 segment E1 leży na skrzyżowaniu COM1 oraz SEG3. Spróbuj zmodyfikować testbench w taki sposób, by zaczernić inne segmenty. Przeprowadź ponowną symulację i załaduj ponownie dane do GTKWave, używając przycisku **Reload** (pierwszy od prawej w górnym pasku narzędzi).

Moduł top

Pozostaje już tylko sporządzić moduł nadrzędny **top**. Jego kod widzimy na **listingu 6**. W module **top** umieścimy instancję generatora sygnału zegarowego, sterownika wyświetlacza oraz prosty licznik 16-bitowy tylko po to, by móc pokazywać na wyświetlaczu LCD jakieś zmieniające się liczby. Testowy licznik będzie inkrementowany co 0,1 sekundy za pomocą modułu **StrobeGenerator**, który wielokrotnie wykorzystywaliśmy w innych aplikacjach. Licznik ma 16 bitów, więc na każdą cyfrę przypadają 4 bity, w których zapisana jest binarnie wartość szesnastkowa od 0 do F. Trzeba ją będzie przekształcić na kod segmentowy za pomocą czterech dekoderek **Decoder7seg**, które opracowaliśmy w odcinku na temat wyświetlacza LED.

Na liście portów modułu **top** mamy jedynie wejście resetujące oraz wyjścia elektrod COM (linia 1) i elektrod SEG (linia 2), które należy podłączyć do filtrów RC, które przekształcają sygnał PWM na sygnał analogowy, wymagany przez wyświetlacz.

Używamy wbudowanego generatora sygnału zegarowego **OSCH**, pracującego z częstotliwością 14 MHz (linia 3). Był już omawiany w poprzednich odcinkach kursu, więc nie będziemy go tutaj omawiać jeszcze raz.

Następnie mamy instancję modułu **StrobeGenerator**, która została skonfigurowana tak, by dawać sygnał do inkrementacji testowego licznika co 100000 mikrosekund, czyli 0,1 sekundy (linia 5). Cyklicznie, co taki czas, generator będzie ustawiał swoje wyjście

Listing 4. Kod pliku **lcd.bat**

```
@echo off
cd impl1
cd source
iverilog -o lcd.o lcd.v lcd_tb.v lcd_pwm.v strobe_generator.v
vvp lcd.o
del lcd.o
pause
```

w stan wysoki na jeden takt zegara. Wyjście generatora jest połączone zmienną wire **CountEnable_w** (linia 4) z warunkiem if, który sprawdzany jest w linii 7 – kiedy warunek jest prawdziwy, następuje zwiększenie testowego licznika **Counter_r** o 1.

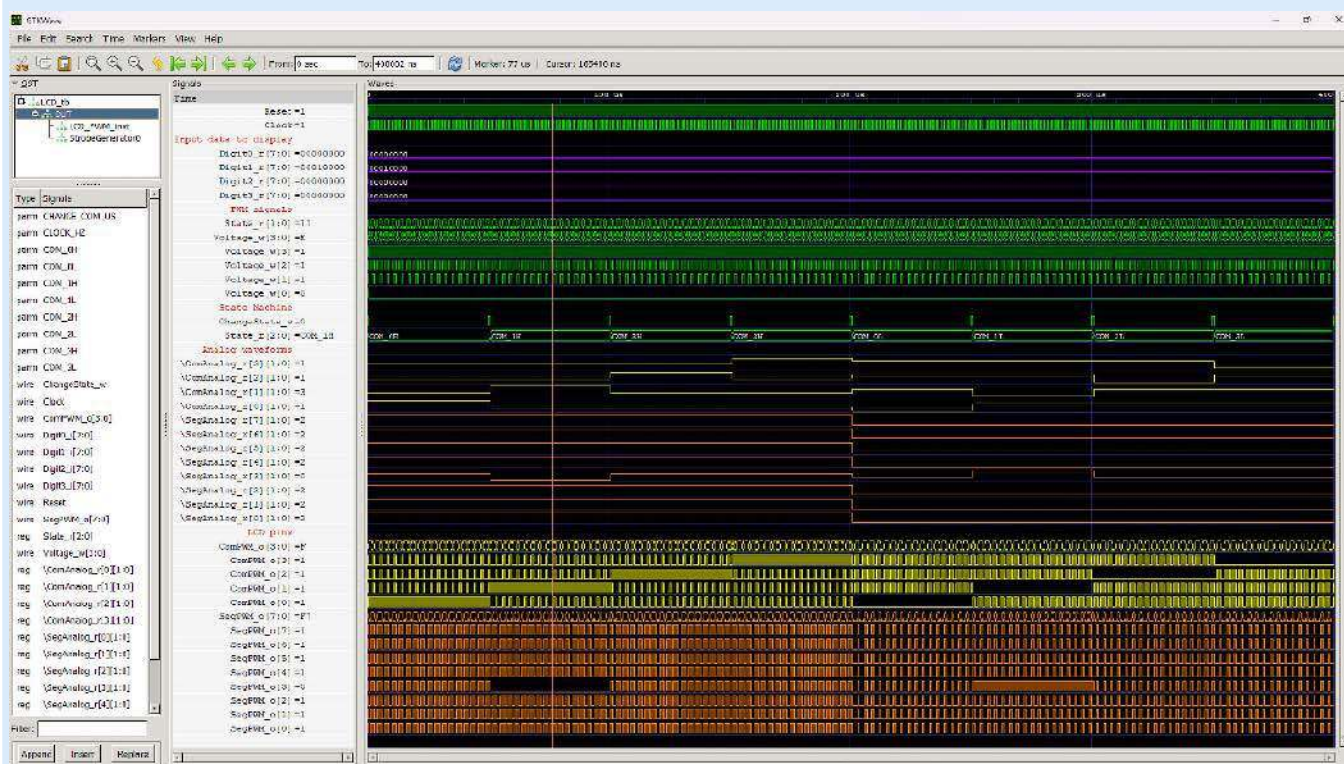
Dalej widzimy cztery instancje dekoderek, przekształcających 4-bitową liczbę binarną na 7-bitowy kod wyświetlacza 7-segmentowego. Na wejściu każdego z dekoderek mamy 4-bitowe „fragmenty” 16-bitowego licznika **Counter_r**, a ich wyjścia są podłączone do wejść danych sterownika wyświetlacza poprzez 7-bitowe zmienne **SegmentsX_w** typu wire.

W linii 10 i kilku kolejnych widzimy wejścia sterownika wyświetlacza. Jednak są to wejścia 8-bitowe, które przyjmują informacje o segmentach w kolejności PGFEDCBA. Segment P to kropka lub dwukropek. Dekodery 7-segmentowe dostarczają tylko informacji o segmentach od A do G. Z tego powodu za pomocą operatora sklejanego { } do danych pochodzących z dekoderek dokleiliśmy zero na początku, aby nie wyświetlać kropek ani dwukropków. Zmienić te zera na jedynki, jeżeli chcesz zaczernić kropki albo dwukropek.

W linii 9 ustawiamy czas trwania każdego z ośmiu stanów wyświetlacza. W naszym przykładzie ten czas jest ustawiony na 5 ms. Spróbuj zmienić tę liczbę na dużo większą, aby dało się zaobserwować, w jaki sposób sterownik załącza kolejne segmenty.

Listing 5. Kod pliku **lcd-state.gtkw**

```
0 COM_0H
1 COM_1H
2 COM_2H
3 COM_3H
4 COM_0L
5 COM_1L
6 COM_2L
7 COM_3L
```



Rysunek 7. Przeglądarka GTKWave z efektami symulacji

Listing 6. Kod pliku top.v

```
// Plik top.v

`default_nettype none
module top(
    input wire      Reset,
    output wire [3:0] ComPWM_o, // 1
    output wire [7:0] SegPWM_o // 2
);

// Generator sygnału zegarowego
parameter CLOCK_HZ = 14_000_000;
wire Clock;
OSCH #(
    .NOM_FREQ("14.00") // 3
) OSCH_inst(
    .STDBY(1'b0),
    .OSC(Clock),
    .SEDSTDBY()
);

// Generator impulsów inkrementujących testowy licznik
wire CountEnable_w; // 4
StrobeGenerator #(
    .CLOCK_HZ(CLOCK_HZ),
    .PERIOD_US(100_000) // 5
) CountEnable_m(
    .Clock(Clock),
    .Reset(Reset),
    .Strobe_o(CountEnable_w)
);

// Testowy licznik
reg [15:0] Counter_r; // 6
always @(posedge Clock, negedge Reset) begin
    if(!Reset)
        Counter_r <= 0;
    else if(CountEnable_w) // 7
        Counter_r <= Counter_r + 1'b1;
end

// Dekoder wyświetlacza - cyfra 0
wire [6:0] Segments0_w;
Decoder7seg Decoder0(
    .Enable_i(1'b1),
    .Data_i(Counter_r[3:0]),
    .Segments_o(Segments0_w)
);

// Dekoder wyświetlacza - cyfra 1
wire [6:0] Segments1_w;
Decoder7seg Decoder1(
    .Enable_i(1'b1),
    .Data_i(Counter_r[7:4]),
    .Segments_o(Segments1_w)
);

// Dekoder wyświetlacza - cyfra 2
wire [6:0] Segments2_w;
Decoder7seg Decoder2(
    .Enable_i(1'b1),
    .Data_i(Counter_r[11:8]),
    .Segments_o(Segments2_w)
);

// Dekoder wyświetlacza - cyfra 3
wire [6:0] Segments3_w;
Decoder7seg Decoder3(
    .Enable_i(1'b1),
    .Data_i(Counter_r[15:12]),
    .Segments_o(Segments3_w)
);

// Instancja sterownika wyświetlacza LCD
LCD #(
    .CLOCK_HZ(CLOCK_HZ),
    .CHANGE_COM_US(5000) // 9
) LCD_inst(
    .Clock(Clock),
    .Reset(Reset),
    .Digit3_i({1'b0, Segments3_w}), // 10
    .Digit2_i({1'b0, Segments2_w}),
    .Digit1_i({1'b0, Segments1_w}),
    .Digit0_i({1'b0, Segments0_w}),
    .ComPWM_o(ComPWM_o),
    .SegPWM_o(SegPWM_o)
);

endmodule
`default_nettype wire
```

Pozostaje już tylko zsyntezować, przypisać piny w Spreadsheet według numeracji, jaką pokazano na rysunku 5, wgrać do FPGA i zobaczyć, jak działa wyświetlacz LCD.

W następnym odcinku nauczymy się generować dźwięki, a żeby nie było monotoni – poznamy, jak działają bloki pamięci EBR wbudowane w MachXO2, w których będą zapisane informacje

o częstotliwościach dźwięków oraz jak długo mają trwać. Zbudujemy prosty odtwarzacz muzyczny o możliwościach kompozytora dzwonków z czasów Nokii 3310.

Dominik Bieczyński
leonow32@gmail.com

REKLAMA

m.technik

Ciekawi świata są zawsze młodzi

w prezencie na każdą okazję
przejrzysz i kupisz na
www.ulubionykiosk.pl



koktajl niusów



Rakietą z napędem hybrydowym z Politechniki Warszawskiej (PW)

17 września 2023 roku o godzinie 16.55 odbył się lot pierwszej w historii Politechniki Warszawskiej (PW) rakietą z napędem hybrydowym. Twardowski – bo tak nazywa się rakietę, stanowi dzieło Studenckiego Koła Astronautycznego PW. Celem, dla którego zbudowano raketę Twardowski, było sprawdzenie nowego rodzaju napędu raketowego, tj. silnika hybrydowego. Podczas lotu, który odbył się na terenie poligonu Lipa, umieszczono na pokładzie rakiet 4 minisatelity, tzw. CanSaty. Rakietę zatankowano zmniejszoną z powodów bezpieczeństwa ilością utleniacza. Silnik rakiet Twardowski działa na podtlenek azotu (N_2O), który jest utleniaczem oraz HTPB, który jest paliwem. Opracowana rakietę wzniosła się na wysokość 1,4 km, po czym nastąpiło wyzwolenie spadochronu spowalniającego. Choć główny spadochron się nie otworzył, to udało się odzyskać i raketę, i CanSaty.

Nad konstrukcją rakiet Twardowski pracowano nieustannie od 2019 roku. Przez ten czas w projekt rakiet włączyło się ponad 70 osób. Ostatnie 2 lata upłynęły głównie na pracy w warsztacie, w którym doprowadzano raketę do gotowości lotnej. Aktualnie trwają prace nad kontynuacją projektu w postaci nowej rakiet, która nosi nazwę Twardovsky 2. Jej parametry znacząco się różnią od poprzedniej wersji – jej docelowy pułap wynosi już 3 km, lecz imponujące 9,144 km, tj. 30 tysięcy stóp. Z tego powodu rakiet będzie znacznie większa niż rakiet Twardovsky.

<https://shorturl.at/sMU28>

VIGI C540V – zewnętrzna, dwuobiektywowa, obrotowa kamera sieciowa firmy TP-Link

Kamera jest przeznaczona do zastosowań zewnętrznych – ma klasę szczelności IP66, zapewnia odporność na trudne warunki atmosferyczne i stabilne działanie na zewnątrz budynku. Dzięki obiektywowi o rozdzielczości 4 Mpx oraz ogniskowej od 4 do 12 mm rejestruje najmniejsze detale. Czuły przetwornik, 4 diody LED i światło punktowe gwarantują w pełni kolorowy obraz także w nocy. Kamera VIGI C540V uwzględnia opcje tworzenia tras patrolu, tj. nieustannego poruszania się między wyznaczonymi



punktami. Może obracać się i pochylać w tempie 44°/s, co pozwala obserwować ogromny obszar. Dodatkowo oferuje zoom 3×, mikrofon, głośnik oraz alarm dźwiękowy i świetlny.

Po ustanowieniu obszarów, kamera wykrywa wtargnięcia na teren, przekroczenie linii, wejście do strefy i jej opuszczenie czy też pozostawienie przedmiotu lub jego zabranie. Wykrywa też osoby, które w podejrzany sposób krążą po obszarze przez określony czas, oraz zmianę sceny, kiedy ktoś np. zasłoni kamerę, żeby uniemożliwić nagrywanie wideo. Użytkownik zyskuje powiadomienie push, kiedy zostanie wykryte którekolwiek z tych zdarzeń. Nagrania z kamery mogą być rejestrowane na rejestratorze sieciowym NVR lub na kartach microSD – do 256 GB. Kamera VIGI C540V korzysta z kompresji H.265+, co minimalizuje obciążenie sieci oraz obniża koszty monitoringu bez utraty jakości obrazu.

Tak jak pozostałe urządzenia z rodziny VIGI, kamera jest zgodna ze standardem ONVIF, dzięki czemu współgra z kamerami i rejestratorami wielu producentów. Za pomocą aplikacji VIGI na urządzeniach przenośnych z systemem iOS lub Android produktami z tej serii można w prosty i kompleksowy sposób zarządzać z poziomu urządzenia mobilnego z każdego miejsca na świecie. Systemem do monitoringu VIGI można też sterować za pomocą odpowiedniego oprogramowania komputerowego. Producent udziela trzyletniej gwarancji na kamerę VIGI C540V.

<https://shorturl.at/tBJLW>



Najnowszy telewizor LG MAGNIT wynosi domową rozrywkę na całkowicie nowy poziom

Telewizor LG MAGNIT jest ekranem o przekątnej 118" i oferuje rozdzielczość obrazu 4K i rozstaw pikseli 0,68 mm. Jego elegancki design oraz bardzo imponujący rozmiar sprawiają, że stanowi idealny wybór nawet do najbardziej ekskluzywnych wnętrz. Dzięki systemowi webOS zapewnia bogaty wybór treści czy wygodne przesyłanie strumieniowe z szerokiej gamy urządzeń. Telewizor LG MAGNIT jest wyposażony w technologię MicroLED, by dostarczyć obraz rewelacyjnej jakości.

Inteligentny procesor Alpha 9 optymalizuje obraz dla filmów, sportu i gier. Procesor ten stosuje zaawansowane algorytmy deep-learning i obszerną bazę danych wizualnych obejmującą szeroki zakres gatunków. Umożliwia to inteligentną analizę oraz dostosowanie jakości wyświetlania do tego, co ogląda użytkownik. Ma to miejsce przy sterowaniu jasnością za pomocą sztucznej inteligencji (AI). Telewizor LG MAGNIT wyposażono w funkcje, które poprawiają szczegóły twarzy i skalują tekst na ekranie, poprawiając komfort oglądania.

Zapewnia obsługę platformy Smart TV webOS z dostępem do rosnącego katalogu popularnych aplikacji streamingowych. Zawiera 2 głośniki – o mocy 50 W, kanał zwrotny audio (eARC), 4 porty HDMI 2.1, a także zintegrowane interfejsy Bluetooth czy WiSA Ready – do bezprzewodowego dźwięku przestrzennego. Opcja montażu na ścianie lub na podstawie pozwala na wygodną instalację telewizora w wybranym miejscu w domu. Telewizor spełnia standardy certyfikacji mające na celu ograniczenie zmęczenia wzroku – emituje niewielką, na dobrą sprawę, ilość niebieskiego światła.

<https://shorturl.at/epO07>



Opaska telemedyczna SiDLY ze statusem wyrobu medycznego

Firma SiDLY otrzymała certyfikat medyczny klasy IIa na własną opaskę telemedyczną stosowaną w teleopiece. Jest to rewelacyjne potwierdzenie tego, że wytwarzana przez firmę opaska to profesjonalne urządzenie, dzięki któremu można zrealizować rzetelne pomiary parametrów życiowych. To nadzwyczaj kluczowe dla adekwatnej oceny stanu zdrowia pacjenta oraz podjęcia właściwych kroków w sytuacji wystąpienia zagrożenia dla zdrowia i życia. Dzięki zyskaniu przez opaskę telemedyczną SiDLY statusu wyrobu medycznego możliwe jest jej legalne stosowanie w placówkach medycznych oraz szpitalach w zakresie monitorowania zdrowia pacjentów. Wspomniana opaska spełnia restrykcyjne wymogi MDR, umożliwiając podjęcie zasadniczych kroków w sytuacji wystąpienia zagrożenia dla zdrowia i życia. W rezultacie można być pewnym, że produkt spełnia najwyższe standardy jakości i bezpieczeństwa, a zespoły medyczne zyskują dostęp do wiarygodnych danych, na podstawie których można podejmować decyzje i postawić diagnozę medyczną.

Jak wyjaśnia prezes firmy SiDLY, Edyta Kocyk: „Certyfikacja medyczna opaski telemedycznej otwiera duże możliwości dla wszystkich! Od teraz zdrowie każdego użytkownika opaski może być przez całą dobę monitorowane, żeby w szybkim tempie móc zareagować”.

<https://shorturl.at/uxBQ4>



INSTAX Pal firmy FUJIFILM – cyfrowy aparat fotograficzny, który mieści się w dłoni

Firma FUJIFILM z przeogromną dumą ogłasza wprowadzenie INSTAX Pal, pierwszego aparatu cyfrowego mieszczącego się w dłoni. Najnowszy aparat serii INSTAX pozwala z łatwością utrwalac w kadrach wszelkie chwile życia. We współpracy z drukarkami

do smartfonów serii INSTAX Link aparat INSTAX Pal to nowa i ekscytująca metoda fotografowania. Za pośrednictwem interfejsu Bluetooth kompaktowy aparat cyfrowy INSTAX Pal może łatwo przesyłać zdjęcia bezpośrednio na smartfon do odpowiedniej aplikacji INSTAX Pal – kiedy wszystko jest gotowe, użytkownik może wydrukować własne ulubione zdjęcia za pomocą jednego z modeli drukarek do smartfonów z rodziny INSTAX Link.

Zdjęcia wykonane aparatem INSTAX Pal mogą być drukowane w formatach mini, SQUARE i WIDE. Nowy aparat pozwala także tworzyć cyfrowe zdjęcia INSTAX, które można zapisać w smartfonie i udostępnić w mediach społecznościowych. Kieszonkowych rozmiarów aparat cyfrowy pozwala realizować zdjęcia jedną ręką. Szerokokątny obiektyw sprawia, że idealnie nadaje się dla potrzeb robienia zdjęć grupowych. Istnieje również możliwość nagrania unikatowego sygnału audio, który można ustawić jako dźwięk odtwarzany tuż przed zwolnieniem migawki aparatu. Użytkownik może dzięki temu nadać urządzeniu swojego, osobistego charakteru. INSTAX Pal występuje w 5 kolorach: Milky White, Powder Pink, Pistachio Green, Lavender Blue, albo Gem Black.

<https://tiny.pl/cld31>



Okrągły jubileusz polskiej marki GOODRAM

W 2023 roku marka GOODRAM obchodzi okrągłe 20. urodziny. Pochodzi ona z Łazisk Górnych na Śląsku i dostarcza niezawodne oraz cenione przez użytkowników moduły RAM, dyski SSD, karty pamięci, a także nośniki USB. Historia marki GOODRAM sięga 1991 roku, kiedy działalność rozpoczęła firma Wilk Elektronik. Początkowo oferta firmy polegała na dystrybucji modułów pamięci DRAM. Wkrótce jednak rozszerzyła się o karty Flash i pamięci USB, a wraz z rozwojem nowych technologii również o dyski SSD. Jednym z ważniejszych punktów zwrotnych było nawiązanie czy utrzymanie w dalszej kolejności współpracy technologicznej z firmą KIOXIA. To strategiczne partnerstwo umożliwiło polskiej marce wdrożenie szeregu rozwiązań Flash do własnej oferty. Przyrost gamy produktów i ekspansja na kolejne rynki eksportowe wywołały ogromną potrzebę, jeśli chodzi o dalszy rozwój. W 2013 roku firma Wilk Elektronik zainwestowała przeszło milion euro w rozwój linii montażu powierzchniowego SMT, udoskonalając efektywność i jakości pamięci GOODRAM. Obecnie firma kończy zaczęta w 2021 roku rozbudowę jedynej w Europie fabryki pamięci komputerowych. Przytoczona inwestycja objęła swoim zasięgiem rozbudowę infrastruktury firmy: od magazynu, poprzez produkcję czy powierzchnie biurowe, aż po nowe rozwiązania w obszarze technologii i przepływu informacji.

<https://utn.pl/G4Hll>

Kolejne szybkie ładowarki ABB na mapie Warszawy

Przy centrum handlowym Factory Annapol w Warszawie działa nowy hub dla pojazdów elektrycznych. Jest to łącznie aż 8 punktów ładowania na 6 stacjach firmy ABB, w tym jedna o mocy 350 kW. Nowa strefa ładowania pokazuje, że dla kierowców i operatora liczy się nie tylko szybkość czy dostępność. Na znaczeniu zyskuje też wsparcie serwisowe oraz funkcjonalność. Na ogólnodostępnym parkingu znajduje



się 6 stacji ładowania, które pozwalają naładować, w całości, każdy rodzaj samochodu elektrycznego. Kierowcy mogą skorzystać z punktów DC – ze złączami CCS oraz punktów AC – ze złączami Type2. Punkty działają w godzinach otwarcia centrum handlowego i od uruchomienia cieszą się znaczącą popularnością wśród kierowców pojazdów elektrycznych.

Pośród zamocowanych ładowarek najszybsza jest Terra High Power 350 trzeciej generacji. Pozwala ona w ciągu zaledwie 3 minut doładować baterię samochodu na następne 100 km jazdy. To pierwsza w Polsce publiczna stacja ładowania ABB o mocy 350 kW, z kablami schładzanymi cieczą.

Oprócz urządzeń, firma ABB dostarczyła również usługę łączności „Charger connect”, która daje zdalny wgląd w stan ładowarek. W rezultacie ok. 95% usterek udaje się zdiagnozować zdalnie, a 75% z nich można naprawić bez konieczności przyjazdu serwisanta na miejsce. Usługa ta umożliwia aktualizacje oprogramowania, a także skorzystanie z narzędzi bazujących na protokole OCPP (standard dla stacji ładowania).

Jak wyjaśnia związany z biznesem Elektryfikacji ABB w Polsce, Łukasz Będziński: „Dzięki chłodzeniu cieczą, moc przekazywana przez ładowarkę do baterii w samochodzie jest na stałym poziomie przez całą sesję ładowania, nawet przy złych warunkach pogodowych. Ogólnodostępnych stacji ładowania z pewnością będzie przybywać. Dostrzegamy rosnące, zainteresowanie ładowarkami szczególnie w obiektach komercyjnych. To właśnie tam, na parkingach, zatrzymuje się coraz więcej samochodów elektrycznych, które prosto i wygodnie mogłyby podczas postoju ładować swoje baterie”.

<https://utn.pl/3nPx8>



Przetwornik obrazu CMOS firmy Sony o oznaczeniu IMX735

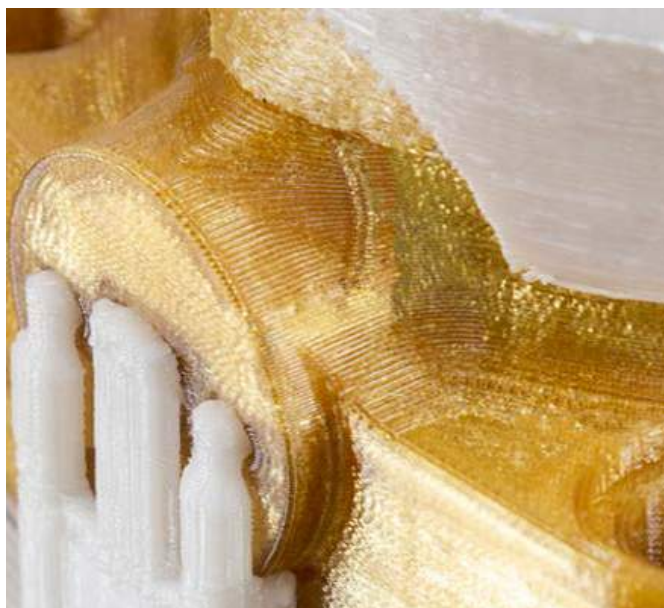
Sony zapowiedziało wprowadzenie na rynek przetwornika obrazu CMOS, o symbolu IMX735. Ten przewidziany dla kamer samochodowych przetwornik cechuje np. rozdzielczość efektywna 17,42 Mpx i ma zagwarantować rozwój systemów kamer samochodowych, zapewniających zaawansowane wykrywanie, a także skuteczność rozpoznawania, przyczyniając się do bezpiecznego, zautomatyzowanego kierowania pojazdem. Aby systemy automatyzacji umożliwiały samoczynne prowadzenie pojazdów, muszą zapewnić rozpoznawanie

otoczenia pojazdu. Rodzi to wzrastający popyt na przetworniki obrazu, które mogą pomóc w spełnieniu opisanych wymagań. Nie inaczej jest, jeżeli chodzi o matrycę z przetwornika IMX735, który pozwala rejestrować obraz wysokiej rozdzielczości – szczególnie z odległości.

Pozostałe przetworniki obrazu CMOS odczytują dane z pikseli pionowo – rząd za rzędem, ten produkt generuje sygnał wyjściowy poziomo, kolumna za kolumną. Ułatwia to synchronizację kamer samochodowych, wyposażonych w ten przetwornik, z lidarami oraz z mechanicznym systemem skanowania. Lepsza synchronizacja znacząco poprawia zdolności zautomatyzowanych systemów kierowania pojazdem do wykrywania oraz rozpoznawania.

Opracowana przez Sony struktura pikseli w przetworniku obrazu IMX735 zwiększa natężenie oświetlenia. Specjalny sposób ekspozycji zapewnia zakres dynamiczny równy 106 dB – nawet w przypadku przetwarzania HDR (*High Dynamic Range*) i ograniczania migotania diod LED. W trybie priorytetu zakresu dynamicznego wartość ta wzrasta do przeszło 130 dB. To interesujące rozwiązanie przeciwdziała prześwietlaniu fragmentów obrazu przy rejestracji pod światło, co zwiększa dokładność rejestracji obiektów w warunkach sporych różnic jasności – np. w czasie wyjazdu z tunelu.

<https://utn.pl/uolsa>



Odporny na wysokie temperatury filament Z-PEI 1010 firmy Zortrax

Filament Z-PEI 1010 jest o wiele mocniejszy niż filament Z-PEI 9085. Może pracować w temperaturach sięgających 208°C. Jest o wiele twardszy od Z-PEI 9085, co czyni go idealnym wyborem w rozwiązaniach, w których liczy się sztywność. Filament Z-PEI 1010 jest odporny na wiele czynników chemicznych, wliczając w to węglowodory halogenowane, alkohole oraz roztwory wodne. Dzięki sporej odporności chemicznej, Z-PEI 1010 sprawdza się w druku elementów rur czy uszczelnień, które nie wchodzi w reakcję ze związkami, z którymi mają kontakt. Zwłaszcza węglowodory halogenowane znajdują zastosowanie w systemach chłodzących, grzewczych oraz urządzeniach przeznaczonych do oczyszczania metali. Dlatego Z-PEI 1010 sprawdza się m.in. w przemyśle chemicznym, motoryzacyjnym i wydobywczym. Filament jest kompatybilny z drukarką 3D Zortrax Endureal i może być użyty w trybie podwójnej ekstruzji z wyłamywanymi podporami wykonanymi z materiału Z-SUPPORT HT. Filament Z-PEI 1010 jest polecany do druku zamienników podzespołów aluminiowych.

<https://tiny.pl/clfhc>

Jakub Tyburski
jakub.tyburski@elportal.pl

Uniwersalny przedwzmacniacz, część 2

Popularne wzmacniacze zintegrowane są zbudowane ze wzmacniaczy mocy i kompletnych układów przedwzmacniaczy. Od dłuższego czasu są również oferowane oddzielne wzmacniacze mocy także w konfiguracji dual mono. Takie wzmacniacze mocy potrzebują do pracy osobnego przedwzmacniacza z układami selektora wejść, regulatorami barwy i poziomu sygnału (siły głosu). Kontynuujemy opis tej niezwykle ciekawej konstrukcji. W drugiej części omówimy budowę bloku przetwornika PCM1794A oraz sterownika mikroprocesorowego.

RX Ewa – odbiornik początkującego nastuchowca

Odbiornik powstał z myślą o początkujących nastuchowcach, którzy pragną zbudować bardzo prosty, tani i o niewielkich wymiarach odbiornik, aby zapoznać się z pracą krótkofalowców na jednym z podstawowych pasm częstotliwości amatorskich 7 MHz, potocznie nazywanych czterdziestką (40 m). Układ został ograniczony do niezbędnego minimum – składa się z popularnych podzespołów i nie wymaga nawijania cewek.

Minimoduł z mikroprocesorem NXP LPC865

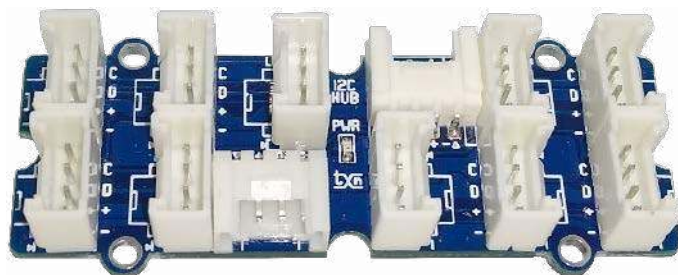
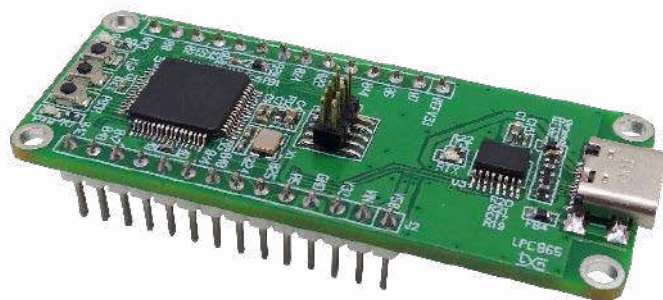
Moduł z najnowszym procesorem NXP typu LPC865M201JBD64 może być ciekawą alternatywą dla popularnych płytek z procesorami AVR czy STM, szczególnie wtedy, gdy nie mamy zbyt wiele miejsca, a nie chcemy rezygnować z wydajności i dobrego wyposażenia procesora. LPC865 to najbardziej rozbudowane układy z rodziny LPC86x M0+, będące kontynuacją serii LPC84x. Wybrany LPC865 z rdzeniem Cortex-M0+, pracującym z zegarem do 60 MHz, ma 64 kB Flash, 8 kB RAM i elastycznie konfigurowane do 54 wyprowadzeń GPIO (obudowa LQFP64).

Aktywny hub I²C

Gdy podczas prototypowania systemu liczba urządzeń na magistrali I²C zwiększa się, może pojawić się problem z niezawodnością transmisji. Pojemności obciążenia linii danych i zegara mogą znacząco zwiększyć czasy narastania sygnału, powodując niepoprawne interpretowanie stanu wysokiego. Skutkiem pogorszenia jakości sygnału może zarządzić akcelerator I²C w postaci układu LTC4311 Analog Devices. Opisany aktywny hub I²C dzięki wyposażeniu w aż 12 złączy w standardzie Grove oraz cztery opcjonalne złącza typu QWIIC ułatwia rozbudowę magistrali i porządkuje okablowanie modułów.

Tematy wiodące w EP 12/2023:

- Systemy akwizycji danych
- Automatyka budynkowa w praktyce



Wykaz firm ogłaszających się w tym numerze „Elektroniki Praktycznej”.

AKSOTRONIK.....	35
ARMEL	13
BL ELEKTRONIK	54, 55
BORNICO.....	5
COMPUTER CONTROLS.....	9
CRC INDUSTRIES EUROPE.....	49
EAE ELEKTRONIK	86
ELMAX.....	79
EUROCIRCUITS.....	83, 90
FERYSSTER	13, 107
GAMMA	13
HAMMOND	7
HATRON	77
LASTENIC LASER.....	69
MICROCHIP	56, 128
PCBWAY	94, 95
PHOENIX CONTACT.....	46, 47
PIEKARZ	13
SATLAND PROTOTYPE.....	75
SECURUS.....	84
WÜRTH.....	82

Miesięcznik „Elektronika Praktyczna” (12 numerów w roku) jest wydawany przez AVT-Korporacja Sp. z o.o. we współpracy z wieloma redakcjami zagranicznymi.



Wydawnictwo:
AVT-Korporacja Sp. z o.o.
03-197 Warszawa, ul. Leszczynowa 11
tel. 22 257 84 99, e-mail: avt@avt.pl

Wydawca:
Wiesław Marciniak

Adres redakcji:
03-197 Warszawa, ul. Leszczynowa 11
e-mail: redakcja@ep.com.pl, www.ep.com.pl

Redaktor Naczelny:
Damian Sosnowski

**Redaktor Programowy,
Przewodniczący Rady Programowej:**
Piotr Zbysiński

Menedżer Magazynu:
Katarzyna Gugala

Szef Pracowni Konstrukcyjnej:
Jakub Sobański

Zespół marketingu i reklamy:

Katarzyna Gugala, tel. 22 257 84 64
Bożena Krzykawska, tel. 22 257 84 42
Grzegorz Krzykowski, tel. 22 257 84 60

Stali współpracownicy:

Lucjan Bryndza, Nikodem Czechowski, Jarosław Doliński,
Andrzej Gawryluk, Krzysztof Górski, Tomasz Jabłoński,
Henryk Kowalski, Rafał Kozik, Michał Kurzela, Przemysław
Musz, Szymon Panecki, Sławomir Skrzyński, Ryszard
Szymaniak, Adam Tatuś, Jakub Tyburski, Robert Wołgajew

Uwaga!

Kontakt z wymienionymi osobami jest możliwy via e-mail,
według schematu: imię.nazwisko@ep.com.pl

DTP i okładka:

MAD Sp. z o.o.

Redakcja strony internetowej www.ep.com.pl

MAD Sp. z o.o.

Prenumerata w Wydawnictwie AVT

www.ulubionykiosk.pl lub tel. 22 257 84 22
(godz. 10.00–14.00)
e-mail: prenumerata@avt.pl

Prenumerata w RUCH S.A.

www.prenumerata.ruch.com.pl
lub tel. 801 800 803, 22 717 59 59
e-mail: prenumerata@ruch.com.pl



Wydawnictwo
AVT-Korporacja Sp. z o.o.
należy do Izby Wydawców Prasy

Copyright AVT-Korporacja Sp. z o.o.

03-197 Warszawa, ul. Leszczynowa 11

Projekty publikowane w „Elektronice Praktycznej” mogą być wykorzystywane wyłącznie do własnych potrzeb. Korzystanie z tych projektów do innych celów, zwłaszcza do działalności zarobkowej, wymaga zgody redakcji „Elektroniki Praktycznej”. Przedruk oraz umieszczanie na stronach internetowych całości lub fragmentów publikacji zamieszczanych w „Elektronice Praktycznej” jest dozwolone wyłącznie po uzyskaniu zgody redakcji. Redakcja nie odpowiada za treść reklam i ogłoszeń zamieszczanych w „Elektronice Praktycznej”.



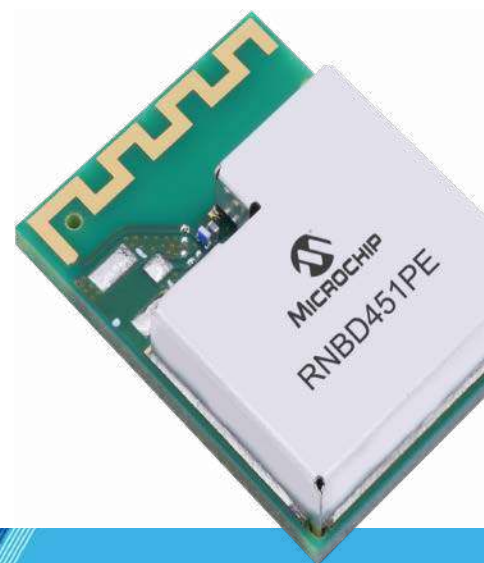


Bez wysiłku dodaj Bluetooth® LE 5.2 z RNBD451

Dodanie Bluetooth® Low Energy (LE) 5.2 do Twojej aplikacji nigdy nie było łatwiejsze niż w przypadku RNBD451. Po prostu włącz i gotowe. Dodaj możliwość bezprzewodowej aktualizacji lub utwórz interfejs dla swojego urządzenia na smartfonie. Możliwości są nieograniczone, a jednocześnie niewymagające dużo pracy.

Kluczowe właściwości

- Gotowa do pracy część RF z optymalnie dobranymi parametrami, nie jest wymagana żadna specjalistyczna wiedza w zakresie układów radiowych
- Nie ma konieczności tworzenia oprogramowania do komunikacji bezprzewodowej, wystarczy wysłać polecenia przez UART
- Globalna certyfikacja oszczędza zarówno czas jak i pieniądze
- Komunikacja Bluetooth® LE o dużym zasięgu z kodowanym PHY
- Funkcja Multi-link umożliwia jednocześnie podłączenie do sześciu urządzeń
- Bezpłatny kod źródłowy aplikacji mobilnej dostępny dla systemów iOS® i Android™



microchip.com/rnbd451pe



eprasa.pl a2e8019fb1

Nazwa i logo Microchip oraz logo Microchip są zastrzeżonymi znakami towarowymi firmy Microchip Technology Incorporated w USA i innych krajach. Wszystkie pozostałe znaki towarowe są własnością ich zarejestrowanych właścicieli.
© 2023 Microchip Technology Inc.
Wszelkie prawa zastrzeżone. MEC2515A-POL-10-23