

ELEKTRONIKA PRAKTYCZNA

EP.com.pl

● Międzynarodowy magazyn elektroników konstruktorów ● październik ● 10/2024 ●

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

- Bardzo prosty odbiornik UKF FM • argbController
- Symulator dźwięku bijącego serca • Miniaturowy zasilacz napięć ustalonych

podzespoły, sprzęt, aplikacje

- Palniki i lutownice gazowe w ofercie TME • Szablony SMT, co warto wiedzieć • Testery funkcjonalne FCT – niezawodne rozwiązania w testowaniu elektroniki
- Udana prototypy sposobem na osiągnięcie celów w zaawansowanych projektach • Sieć Thread w automatyzacji budynków

tutoriale

- Parametry optyczne źródeł światła LED • Oświetlenie LED w aplikacjach konsumenckich i profesjonalnych
- Moduły LED COB – rodzaje i aspekty aplikacyjne
- Produkcja płytek drukowanych – przewodnik w pigułce • Internet Rzeczy w pomiarach środowiskowych. Czujnik UV z fotodiodą GUA-S12SD • Prototypowe sztuczki i kruczki
- Amatorskie pomiary wzmacniaczy audio • Wejścia różnicowe w urządzeniach audio • Oszczędzanie energii w teorii i w praktyce

kursy

- Kurs FPGA Lattice. Miernik częstotliwości • Kurs Nordic nRF z BT. Bluetooth LE

ŹRÓDŁA ŚWIATŁA LED

TEMAT NUMERU



PŁYTKI PCB I PROTOTYPY



-20%
NA START
181,40 zł

-30%
po pierwszym roku
prenumeraty
158,80 zł

-40%
po drugim roku
prenumeraty
136,10 zł

-50%
po trzecim roku
nieprzerwanej prenumeraty
113,40 zł

Odkryj korzyści z **prenumeraty drukowanej** – większe oszczędności z każdym rokiem!

Rozpocznij swoją przygodę z *Elektroniką Praktyczną*. Decydując się teraz na roczną prenumeratę drukowaną, otrzymasz nie tylko dostęp do najnowszych wydań, ale i **znakomity start dzięki zniżce 20%** na pierwsze zamówienie!

Prenumerata to nie tylko wygoda dostępu do treści, ale także sposób na znaczące oszczędności. Dołącz do grona naszych stałych czytelników i ciesz się coraz lepszymi warunkami.

Im dłużej jesteś z nami, tym więcej oszczędzasz:

- po roku nieprzerwanej prenumeraty zapewnimy Ci **30% rabatu** na kolejny rok,
- po dwóch latach wierności zaoferujemy **40% rabatu**,
- po trzech latach lojalności osiągniesz **najwyższy poziom rabatu – 50%!**

Jak otrzymać rabat za lojalność?

Zaloguj się na swoje konto prenumeratora na www.UlubionyKiosk.pl i zamów prenumeratę, korzystając z przycisku PRZEDŁUŻ w zakładce „Prenumeraty”.

Przeglądaj wcześniej, płać mniej – postaw na **e-prenumeratę!**

Wybierz prenumeratę cyfrową PDF i ciesz się dostępem do czasopisma nawet 7 dni przed oficjalną premierą w kioskach. Oszczędzaj czas i pieniądze – skorzystaj z **rabatu 30%** na roczną e-prenumeratę w cenie 126,90 zł.

Dodatkowa oferta dla prenumeratorów wersji drukowanej: jeśli już subskrybujesz wersję papierową, możesz dokupić równoległe e-wydania w cenie 36,20 zł/rok – z **niesamowitym rabatem 80%**.

Zyskaj nieograniczony dostęp do zasobów dla pasjonatów elektroniki!

Tylko prenumeratorzy mają pełny dostęp do:

- artykułów przed ich publikacją w *Elektronice Praktycznej* na www.ep.com.pl – EP W TOKU
- materiałów dodatkowych (takich jak pliki źródłowe projektów) na www.UlubionyKiosk.pl/media

Zamów prenumeratę drukowaną lub e-prenumeratę na www.UlubionyKiosk.pl lub przez przelew na konto Wydawnictwa AVT, a po zaksięgowaniu wpłaty wyślemy Ci mailowo kod dostępu do portalu.



Zacznij korzystać z pełnych zasobów już dziś!

Elektronika w służbie społeczeństwa

W chwili gdy piszę te słowa, przez południowo-zachodnią część naszego kraju przetacza się powódź. Największa od pamiętnego roku 1997. Niszcząca siła niewyobrażalnie wielkich mas wody zabiera dobytek tysięcy osób, porywa samochody, przemienia obraz niewielkich miejscowości w postapokaliptyczny chaos. Trudno się zatem dziwić, że mieszkańcy obszarów położonych w niższych partiach biegu rzek w napięciu oczekują dalszego rozwoju sytuacji, śledzą doniesienia lokalnych mediów, a niektórzy – niestety wbrew zaleceniom służb – udają się na wały przeciwpowodziowe, by na własne oczy zobaczyć zbliżający się żywioł.

Zaraz, zaraz... ale jaki związek mają katastrofalne wydarzenia września 2024 z elektroniką? Wbrew pozorom – bardzo duży. Najprawdopodobniej jest to bowiem pierwszy w historii Polski (a na pewno pierwszy na tak dużą skalę) kataklizm, w którego opanowywaniu – z oddolnej inicjatywy – tak czynny udział biorą mieszkańcy wyposażeni w najnowocześniejszy sprzęt. Nad głównymi nurtami rzek cały czas krążą dziesiątki prywatnych dronów, których operatorzy co rusz dzielą się najnowszymi zdjęciami i nagraniami ze społecznością zgromadzoną na lokalnych grupach w mediach społecznościowych. I nie chodzi tutaj tylko o zaspokojenie oczywistej ciekawości czy też próbę poradzenia sobie z lękiem przed nieznanym przebiegiem dalszych wydarzeń. Już teraz bowiem pojawiają się koncepcje obywatelskich patroli – niewielkie, amatorskie drony wzbijają się na maksymalną dopuszczalną wysokość, śledzą bieg wypadków i... pozwalają na wczesne wykrywanie zagrażających pobliskim domostwom przecieków. Im szybciej służby i wolontariusze dotrą na miejsce uszkodzenia wału, tym większa szansa na uchronienie go przed zniszczeniem i na zapobieżenie wdarciu się wody w głąb suchego jeszcze lądu. Prywatne, niewielkie drony o masie poniżej 250 g nigdy chyba nie brały jeszcze udziału w tak ważnej akcji, która – podkreślmy to raz jeszcze – ma swoje źródło w całkowicie oddolnej inicjatywie.



Na tym jednak nie kończy się rola elektroniki w ochronie przed wielką wodą. Coraz większe znaczenie zyskują media społecznościowe, umożliwiające szybkie organizowanie pomocy w najbardziej zagrożonych obszarach. Zarządzanie pracą osób, które z potrzeby serca rzucają codzienne zajęcia i w pocie czoła walczą z żywiołem, także odbywa się poprzez social media (podobnie zresztą, jak indywidualne inicjatywy dostarczania posiłków dla wojska, policji oraz wolontariuszy, organizowane przez lokale gastronomiczne). Nie sposób nie wspomnieć także o tych w pełni profesjonalnych zastosowaniach elektroniki, która pracuje na pierwszej linii frontu walki z żywiołem – rozbudowana sieć telemetryczna monitoruje poziom wody w korytach i zbiornikach retencyjnych, co ma fundamentalne znaczenie w procesie decyzyjnym instytucji odpowiedzialnych za regulację rzek oraz służb mundurowych zabezpieczających kluczowe punkty.

Niestety, jak zawsze, jest także druga strona medalu – znacznie bardziej ponura. Te same kanały komunikacji, które odgrywają tak istotną rolę w bieżącym przekazywaniu ważnych danych, są także wykorzystywane przez internetowych trolli siejących dezinformację, bezmyślnie szerzących panikę lub wręcz przeciwnie: świadomie manipulujących faktami w taki sposób, by przekonać mieszkańców, że wszelkie przygotowania i zabezpieczenia są bezcelowym wysiłkiem. Coraz częściej też w mediach społecznościowych do głosu dochodzą pseudoeksperci, którzy – pomimo całkowitego braku przygotowania merytorycznego z dziedziny hydrologii czy zarządzania kryzysowego – próbują szerzyć swoje wydumane „mądrości”, co dodatkowo zwiększa chaos informacyjny. Oficjalne komunikaty giną wśród zalewu niesprawdzonych informacji, propagujących drogą internetowego głuchego telefonu.

Czy mamy możliwość obrony przed fake newsami? Zdecydowanie tak i to na różnych płaszczyznach. Z jednej strony trzeba uświadamiać bliskich i znajomych o istnieniu „kampanii dezinformacyjnej” i przekonywać do korzystania tylko z pewnych źródeł informacji oraz mniejszego zaufania do „mądrości z internetów”. To właśnie do nas – osób świetnie obeznanych z nowymi technologiami – należy to edukacyjne zadanie.

Z drugiej strony należy spodziewać się w najbliższych latach rozwoju technologii mających na celu weryfikację komentarzy, postów i artykułów bazujących na fake newsach. Mieliśmy już zresztą próbkę działania tego typu funkcjonalności podczas niedawnej pandemii COVID-19 – przeglądając informacje na temat samej choroby, metod leczenia czy też zapobiegania (m.in. za pomocą szczepień), natrafialiśmy często na komunikaty przypominające o aktualnych, oficjalnych zaleceniach prewencyjnych. To samo dzieje się nawet w aspektach walki z internetowym hejtem, który dla wielu osób – zwłaszcza tych najmłodszych i nieodpornych na „nowoczesne formy przemocy” – skończył się już tragicznie. Centrum Badań nad Uprzedzeniami UW przeprowadziło niedawno interesujące badanie z użyciem bota AI, który – działając z poziomu konta na portalu reddit – automatycznie odpisywał agresorom na ich komentarze, zwracając uwagę na hejterski charakter wiadomości. Co ciekawe, taka „demaskacja” skutecznie obniżała liczbę nienawistnych interakcji. Sztuczna inteligencja zaczęła więc... uzupełniać braki wychowawcze niektórych internautów!

Jak zawsze efekty korzystania z technologii zależą więc od intencji jej użytkowników. Mam nadzieję, że wszyscy nasi Czytelnicy w codziennym życiu spotykać będą możliwie jak najwięcej pozytywnych przykładów zastosowań IT. I choć świat, w którym technologia działa tylko na korzyść ludzkości, jest utopią – to mimo wszystko warto, by każdy z nas we własnym zakresie starał się zbliżyć do tego ideału. Czego Wam, drodzy Czytelnicy, a także sobie szczerze życzę.

Zapraszam do lektury!

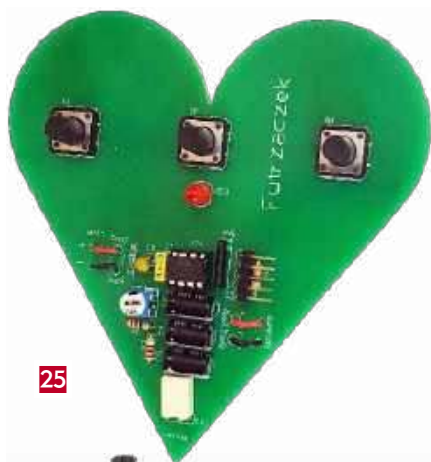
Przemysław Musz



12



17



25



28



64

Nie przeocz

Nowe podzespoły	5
Dodaj do obserwowanych	10
Koktajl niusów	92

Projekty

Bardzo prosty odbiornik UKF FM	12
argbController (2)	17

Miniprojekty

Symulator dźwięku bijącego serca	25
Miniaturowy zasilacz napięć ustalonych	28

Temat numeru: Nowoczesne zasilacze impulsowe

Parametry optyczne źródeł światła LED	30
Oświetlenie LED w aplikacjach konsumenckich i profesjonalnych	39
Moduły LED COB – rodzaje i aspekty aplikacyjne	47

Prezentacje

Palniki i lutownice gazowe w ofercie TME	22
Szablony SMT, co warto wiedzieć	50
Testery funkcjonalne FCT	
– niezawodne rozwiązania w testowaniu elektroniki	53
Udane prototypy sposobem na osiągnięcie celów	
w zaawansowanych projektach	63
Sieć Thread w automatyzacji budynków	90

Elektronika w praktyce

Produkcja płytek drukowanych – przewodnik w pigułce	54
---	----

Moduły w aplikacjach

Internet Rzeczy w pomiarach środowiskowych (10).	
Czujnik UV z fotodiodą GUVA-S12SD	64

Notatnik konstruktora

Prototypowe sztuczki i kruczki	68
Amatorskie pomiary wzmacniaczy audio (4)	70
Wejścia różnicowe w urządzeniach audio	74
Oszczędzanie energii w teorii i w praktyce (4)	76

Kursy

Kurs FPGA Lattice (24). Miernik częstotliwości	81
Kurs Nordic nRF z BT (4). Bluetooth LE	85

Prenumerata	2
Od wydawcy	5
Hity następnego numeru	95

nowe podzespóły

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



2-kanalowe, jednopłytkowe sterowniki bramek modułów IGBT o napięciu do 2,3 kV

Firma Power Integrations wprowadza na rynek nową rodzinę dwukanałowych sterowników bramek modułów IGBT o ratingu napięciowym 1,2...2,3 kV, zrealizowanych na pojedynczej płytce drukowanej. Są to sterowniki plug-and-play współpracujące m.in. z modułami LV100 (Mitsubishi), XHP 2 (Infineon), HPnC (Fuji Electric) oraz ekwiwalentnymi o napięciu blokowania do 2300 V, znajdującymi zastosowanie m.in. w instalacjach fotowoltaicznych, na farmach wiatrowych oraz w systemach przechowywania energii.

Dzięki małym gabarytom nowe sterowniki mogą być zamocowane w obrysie modułu IGBT, co daje projektantom większą swobodę, jeśli chodzi o konstrukcję mechaniczną. Zawierają układ aktywnej kontroli temperatury, współpracujący z wbudowanym w moduł IGBT termistorem NTC – pozwala to na optymalizację konstrukcji termicznej oraz uzyskanie większej o 25...30% mocy wyjściowej przy mniejszej liczbie komponentów, kabli, złączy i barier izolacyjnych. Do standardowego wyposażenia sterowników należą ponadto: układ soft-shutdown, zabezpieczenie przeciwzwarceniowe oraz zabezpieczenia podnapięciowe sekcji pierwotnej i wtórnej.

Obecnie w ofercie Power Integrations dostępne są dwa pierwsze modele: 2SIXT0112T2A0-CM1200DW-34T, przeznaczony

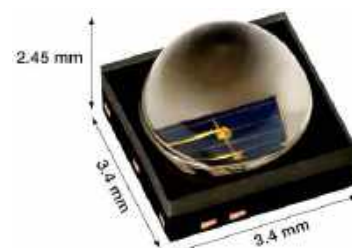
do współpracy z modulem CM1200DW-34T produkcji Mitsubishi oraz 2SIXT0112T2A0-FF1800XTR17T2P5, dostosowany do współpracy z modulem FF1800XTR17T2P5 marki Infineon. Oba charakteryzują się mocą wyjściową 1,0 W/kanał, napięciem włączania/wyłączania bramki równym +15/-10 V, napięciem zasilania 15 V, maksymalną częstotliwością przełączania 25 kHz i zakresem dopuszczalnej temperatury pracy od -40 do +85°C. Zawierają pokrycie konformalne, zabezpieczające przed kondensacją wilgoci.

www.power.com

Emiterzy IR z kwalifikacją AEC-Q102, o małych gabarytach i dużej mocy promieniowania

Oferta emiterów podczerwieni firmy Vishay powiększyła się o 8 nowych modeli z kwalifikacją AEC-Q102, zaprojektowanych do zastosowań w elektronice samochodowej. Zapewniają one odporność na wyładowania ESD do 5 kV, a zakres dopuszczalnej temperatury pracy rozciąga się od -40 do +125°C. W porównaniu z najlepszymi odpowiednikami dostępnymi obecnie na rynku, wyróżniają się zwiększoną o 10% mocą promieniowania przy mniejszej o 20% powierzchni montażowej, co pozwala zmniejszyć liczbę komponentów optoelektronicznych w docelowej aplikacji i ograniczyć koszty.

Poza samochodowymi systemami ADAS oraz rozwiązaniami do monitorowania kierowcy i wnętrza kabiny, nowe emiterzy IR mogą być stosowane również w telewizji przemysłowej. Występują w wersjach o długości fali 850 nm i 940 nm oraz o kącie emisji od $\pm 28^\circ$ do $\pm 75^\circ$. Mogą być sterowane prądem ciągłym o natężeniu do 1,5 A (do 5 A w impulsie). Zapewniają ciągłą moc promieniowania do 2000 mW/sr i szczytową do 6000 mW/sr. Są produkowane



REKLAMA

HAMMOND®
1557 Obudowy poliwęglanu IP68

Dowiedz się więcej:
hammondmfg.com/1557



eusales@hammfg.com • + 44 1256 812812



	Wymiary (mm)	Długość fali (nm)	Moc promieniowania (mW/SR) @		Kąt emisji (°)	Czas narastania (ns)
			I _F =1,5 A	I _F =5 A		
V SMA1094750X02	3,4×3,4×1,5	940	535	1600	±75	10
V SMA1094600X02	3,4×3,4×1,8		750	2300	±60	10
V SMA1094400X02	3,4×3,4×2,45		1525	4620	±40	10
V SMA1094250X02	3,4×3,4×2,9		2000	6000	±28	10
V SMA1085750X02	3,4×3,4×1,5	850	535	1600	±75	13
V SMA1085600X02	3,4×3,4×1,8		750	2300	±60	13
V SMA1085400X02	3,4×3,4×2,45		1525	4620	±40	13
V SMA1085250X02	3,4×3,4×2,9		2000	6000	±28	13

w obudowach SMD o powierzchni 3,4×3,4 mm i małej rezystancji termicznej, wynoszącej – w zależności od modelu – od 5 do 9 K/W.

www.vishay.com

9-megapikselowy czujnik obrazu typu Global Shutter do kamer przemysłowych

Na wystawie Automate Show firma Omnivision zaprezentowała nowy, 9-megapikselowy czujnik obrazu do kamer przemysłowych. OG09A to 1-calowa matryca CMOS BSI (backside illuminated) o rozdzielczości 4096×2160 pikseli @ 60 fps, pracująca z migawką globalną (dobrze sprawdzającą się podczas rejestrowania szybko przemieszczających się obiektów). Może znaleźć zastosowanie w automatyzacji produkcji i inteligentnych systemach transportowych. Na tle innych czujników tej klasy wyróżnia się małymi szumami i dużą sprawnością kwantową. Dzięki znacznej średnicy piksela (3,45 μm) oraz zastosowanym technologiom Nyxel NIR (near-infrared) i DCG HDR (Dual Conversion Gain High Dynamic Range) zapewnia szeroki zakres dynamiczny (76 dB) oraz czytelny obraz, również w miejscach o słabym oświetleniu. Transmisja danych odbywa się przez 16-liniowy interfejs LVDS z szybkością do 1,05 Gbps.

OG09A występuje w wersji kolorowej (ozn. OG09A10) i monochromatycznej (OG09A1B).

Pozostałe cechy:

- 12-bitowy przetwornik A/C,
- wbudowany czujnik temperatury,
- wbudowana pamięć OTP,
- interfejs SPI do programowania wewnętrznych rejestrów,
- automatyczna korekta poziomu czerni,
- obsługa trybów ROI.

www.ovt.com

Switch zabezpieczający do portów USB-C o wydajności prądowej 7 A

AOZ1377DI to nowy typ switcha zabezpieczającego do portów USB-C zgodnych ze specyfikacją Type-C PD 3.0. Może on pracować z maksymalnym napięciem wejściowym 23 V i prądem wyjściowym do 7 A. Zawiera ogranicznik prądowy, programowany zewnętrznym rezystorem, układ miękkiego startu oraz zabezpieczenia: przed odwrotną polaryzacją napięcia, podnapięciowe, nadnapięciowe, termiczne i przeciwzwarciove. Wewnętrzne tranzystory N-MOSFET, połączone w układzie back-to-back, zapewniają bardzo małą rezystancję przewodzenia (typ. 19 mΩ) i szeroki obszar bezpiecznej pracy (SOA), co pozwala na zastosowania również w aplikacjach z dużymi obciążeniami pojemnościowymi. Współczynnik slew-rate można regulować za pomocą zewnętrznego kondensatora.

AOZ1377DI może być przełączany w tryb uśpienia, w którym pobór prądu ulega ograniczeniu do około 3 μA. Występuje w dwóch wariantach: AOZ1377DI-01 – z automatycznym restartem po zadziałaniu któregoś z wbudowanych zabezpieczeń – oraz AOZ1377DI-02,



wymagającym ponownego włączenia układu sygnałem podanym na wejście Enable.

AOZ1377DI spełnia wymogi norm IEC61000-4 w zakresie kompatybilności elektromagnetycznej, a także IEC62368-1:2018 (3rd Edition) w zakresie bezpieczeństwa użytkownika. Jest produkowany w obudowie DFN3x3-10L. Jego ceny hurtowe zaczynają się od 1,36 USD przy zamówieniach 1000 sztuk.

www.aosmd.com



Miniaturowy czujnik CO₂ o poborze prądu 850 μA

STCC4 to jeden z najmniejszych na rynku czujników CO₂, oferowany w obudowie SMD o wymiarach 4×3×1,2 mm. Został on zaprojektowany do wszelkiego typu aplikacji (głównie klimatyzatorów, inteligentnych termostatów i systemów monitorowania jakości powietrza wewnątrz pomieszczeń), w których najważniejszymi kryteriami są gabaryty oraz cena podzespołów. Firma Sensirion zapowiedziała rozpoczęcie masowej produkcji tego układu pod koniec 2024 roku.

STCC4 charakteryzuje się zakresem pomiarowym od 400 do 5000 ppm, dokładnością ±100 ppm i liniowością ±10% MV ppm. Pracuje z napięciem zasilania 3,3 V, pobierając średnio 850 μA prądu, a maksymalnie 5 mA. Komunikuje się z mikroprocesorem za pośrednictwem interfejsu I²C.

www.sensirion.com

Złącza dużej gęstości NovaRay do transmisji sygnałów z szybkością 112 Gbps/kanal

Do oferty firmy Samtec wchodzi złącza dużej gęstości NovaRay o powierzchni o 40% mniejszej od wcześniejszych odpowiedników. Umożliwiają one transmisję danych z szybkością do 112 Gb/s na kanal w najbardziej wymagających zastosowaniach, m.in. w sektorach: półprze-



wodnikowym, medycznym czy wojskowym. Mogą pracować w szerokim zakresie temperatury otoczenia: od -40 do $+125^{\circ}\text{C}$. Wykazują bardzo małe przesłuchy międzykanałowe (<30 dB FEXT, NEXT) w zakresie częstotliwości do 40 GHz. Spełniają wymogi specyfikacji PCIe 6.0/CXL 3.1.

Innowacyjna konstrukcja złączy NovaRay zapewnia doskonałe parametry przy dużej gęstości kanałów. Są to złącza różnicowe z pełnym ekranowaniem i kontaktami o dwóch niezależnych punktach styku. W ramach systemu NovaRay producent dostarcza także gotowe kable zawierające dwa złącza, wraz z opcjonalnym systemem zatraskowym oraz ekranowaniem. Wszystkie wiązki zawierają złącze NovaRay, a także – w zależności od wersji – złącze FQSP, NovaRay I/O, ExaMAX backplane lub NovaRay backplane. Jeśli chodzi o przewody, standardowo dostępny jest wariant 34 AWG EyeSpeed o bardzo małej różnicy czasów propagacji między liniami ($<3,5$ ps/m), a opcjonalnie 34 AWG/100 Ω , 34 AWG/92 Ω i 34 AWG ThinaxTM/92 Ω (o mniejszym o 40% przekroju).

www.samtec.com

Resetowalny bezpiecznik PPTC z kwalifikacją AEC-Q200, o prądzie trzymania równym 3 A

W ofercie resetowalnych bezpieczników PPTC firmy Bourns pojawia się nowy model, stanowiący rozszerzenie serii MF-SMHT. Jego znakiem rozpoznawczym jest prąd trzymania (IHOLD) zwiększony do 3,0 A. Wcześniejsze modele występowały w wersjach o wartości IHOLD do 1,6 A.

Podobnie jak wcześniejsze bezpieczniki MF-SMHT, również nowy model MF-SMHT300 charakteryzuje się maksymalnym dopuszczalnym prądem i napięciem roboczym do (odpowiednio)



	V _{maks.}	I _{maks.}	I _{HOLD} (+23°C)	I _{TRIP} (+23°C)	Rezystancja (+23°C)	Maks. czas zadziałania
MF-SMHT136	16 V	100 A	1,36 A	2,72 A	0,085...0,33 Ω	10 s @ 8 A
MF-SMHT160			1,6 A	3,2 A	0,050...0,15 Ω	10 s @ 8 A
MF-SMHT300			3,0 A	6,0 A	0,024...0,083 Ω	8 s @ 15 A

100 A/16 V oraz uzyskał kwalifikację AEC-Q200. Może pracować w temperaturze otoczenia od -40 do $+125^{\circ}\text{C}$. W zakresie odporności na wibrację spełnia wymogi normy MIL-STD-883C (Method 2007.1, Condition A). Jego żywotność wynosi 1000 h w temperaturze $+125^{\circ}\text{C}$ ($\Delta R = \pm 15\%$).

MF-SMHT300 jest produkowany w obudowie SMD o wymiarach 9,5×6,7×3,0 mm.

www.bourns.com

Wzmacniacz mocy na pasmo 27,5...31 GHz do terminali satelitarnych

Firma CML Micro wprowadza do oferty nowy wzmacniacz mocy CMX90A705 na pasmo Ka, wykonany w technologii GaN SiC 0,15 μm . Jest to wzmacniacz dwustopniowy o maksymalnej mocy wyjściowej +37,4 dBm (5,5 W), pracujący w zakresie częstotliwości od 27,5 do 31 GHz. Oferuje wzmocnienie małosygnałowe na poziomie 16,5 dB, sprawność (PAE) równą 22% @ Psat i współczynnik ACPR lepszy od -28 dBc @ 30 dBm. Pracuje z napięciem zasilania 27,5 V, pobierając 100 mA prądu.

CMX90A705 może odgrywać rolę zarówno stopnia sterującego, jak i końcowego wzmacniacza mocy w terminalach do komunikacji satelitarnej. Został zaprojektowany z myślą o łatwej integracji



REKLAMA

+

W ZESTAWIE TANIEJ

www.ccontrols.pl

Tel: +48 (33) 485 94 90
E-mail: info@ccontrols.pl

w systemie. Zawiera porty wejściowe i wyjściowe dopasowane do impedancji 50 Ω oraz kondensatory blokujące DC. Jest produkowany w obudowie QFN o powierzchni 4x4 mm.

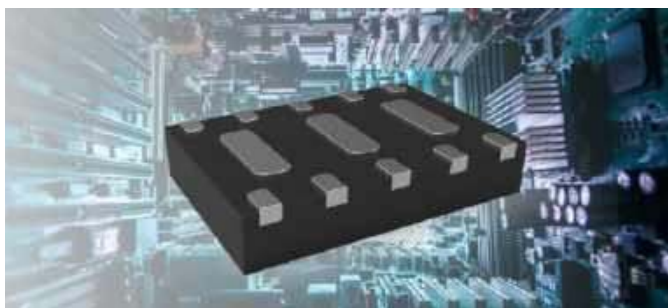
www.cmlmicro.com

Pierwsze 1- i 2-megabitowe pamięci nieulotne F-RAM do pracy w przestrzeni kosmicznej

Infineon powiększa ofertę pamięci przystosowanych do pracy w przestrzeni kosmicznej o pierwsze 1- i 2-megabitowe pamięci nieulotne F-RAM z interfejsem równoległym, wykazujące podwyższoną odporność na promieniowanie jonizujące. Układy te są oferowane w ceramicznych obudowach TSOP-44 i QML-V. Mogą pracować w militarnym zakresie temperatury otoczenia od -55 do +125°C. Są odporne na całkowitą dawkę napromieniowania >150 krad (SI) oraz impulsy SEL o energii >96 MeV·cm²/mg @ +115°C. Producent deklaruje ich niezawodność na poziomie 1013 operacji odczytu/zapisu oraz 120-letni czas retencji danych w temperaturze +85°C.

CYRS15B101N i CYRS15B102N mogą być używane do przechowywania wyników pomiarów z czujników, danych kalibracyjnych, kluczy szyfrujących oraz kodu sterującego mikrokontrolerów. Charakteryzują się pojemnością odpowiednio: 1 Mb (128 kx8) i 2 Mb (128 kx16). Pracują z napięciem zasilania 2,0...3,6 V przy maksymalnym poborze prądu 20 mA w stanie aktywnym, 700 μ A w trybie standby i 20 μ A w trybie uśpienia. Zapewniają odporność na wyładowania ESD do 2 kV (HBM) i 500 V (CDM).

www.infineon.com



4-kanalowy układ zabezpieczający przed wyładowaniami ESD do ± 30 kV

ESDLC3304P9 to układ realizujący zabezpieczenie 4 par linii sygnałowych przed wyładowaniami ESD, dostępny w obudowie DFN3020-10. Charakteryzuje się on maksymalnym napięciem wstecznym VRWM równym 3,3 V przy prądzie upływu do 0,5 μ A, małym napięciem ograniczenia (clamping voltage), nieprzekraczającym 5,5 V dla IPP=1 A oraz typową pojemnością wewnętrzną 1,7 pF, pozwalającą na zastosowania na szybkich liniach transmisyjnych.

ESDLC3304P9 zapewnia ochronę przed wyładowaniami ESD (przenoszonymi przez powietrze i dotyk) do ± 30 kV (zgodnie z IEC 61000-4-2) oraz wytrzymuje impulsy prądowe do 40 A (8/20 μ s), zgodnie z IEC 61000-4-5. Jego dodatkową zaletą stanowi szeroki zakres dopuszczalnej temperatury pracy złącza: od -55 do +125°C, umożliwiając pracę w środowiskach przemysłowych.

www.mccsemi.com

Nowy system złączy krawędziowych kabel-płytki niewymagający lutowania

Nowe złącza krawędziowe kabel-płytki serii 9169-000 są wyposażone w styki kompresyjne z miedzi berylowej oraz styki IDC (z przemieszczeniem izolacji), ułatwiające realizację niezawodnych



połączeń o dużej integralności sygnału – bez konieczności lutowania – w systemach oświetleniowych i przemysłowych oraz w elektronice użytkowej. Złącza te zapewniają doskonałe parametry elektryczne i mechaniczne w trudnych warunkach pracy, m.in. w motoryzacji czy transporcie. Zawierają dwa pola zatraskowe, zapobiegające przypadkowemu rozłączeniu oraz dodatkowe wycięcie, zwiększające stabilność mechaniczną i eliminujące ryzyko przypadkowego odwrócenia polaryzacji podczas montażu.

W ramach serii 9169-000 dostępne są warianty zawierające od 2 do 10 kontaktów z pokryciem cynowym lub pozłacanym, rozdzielonych co 2,5 mm. Umożliwiają montaż przewodów litych lub linkowych o średnicy od 0,32 do 0,64 mm. Wyróżniają się szerokim zakresem dopuszczalnej temperatury pracy (od -40 do +130°C) i niepalną konstrukcją (UL94 V-0). Ich napięcie robocze wynosi 100 VAC rms, a dopuszczalny prąd przewodzenia to 3 A w przypadku wersji zawierających od 2 do 6 kontaktów i 2 A dla wersji o 7...10 kontaktach.

Złącza z serii 9169-000 są dostępne w 8 wariantach kolorystycznych: białym, czarnym, pomarańczowym, zielonym, czerwonym, żółtym, niebieskim i brązowym.

www.kyocera-avx.com

4-parametrowy czujnik jakości powietrza wewnątrz pomieszczeń

BME690 to kolejny czujnik firmy Bosch Sensortec, przeznaczony do badania jakości powietrza wewnątrz pomieszczeń. Jest czujnikiem 4-parametrowym, zrealizowanym w technologii MEMS, mierzącym wilgotność, temperaturę, ciśnienie atmosferyczne oraz zawartość lotnych związków organicznych (VOC) i lotnych związków siarki (VSC). W porównaniu z wcześniejszymi odpowiednikami, charakteryzuje się mniejszym o połowę poborem mocy, co – w połączeniu z małymi gabarytami (3x3x0,93 mm) – pozwala na stosowanie go w urządzeniach bateryjnych.

BME690 stanowi unowocześnioną wersję wcześniejszych czujników BME688 i BME680, oferującą dodatkowo mechanizmy sztucznej inteligencji, do których dostęp zapewnia dostarczana przez producenta aplikacja BME AI-Studio. Jest zamykany w identycznej obudowie, co ułatwia ewentualny upgrade istniejących urządzeń. Pracuje z napięciem zasilania od 1,71 do 3,6 V, pobierając od 2,1 μ A do 3,9 mA prądu – w zależności od trybu pracy i liczby mierzonych parametrów. Jest zgodny z wymogami standardów WELL i RESET, dotyczących monitorowania jakości powietrza w pomieszczeniach. Pracuje w zakresie ciśnienia 300...1100 hPa, wilgotności względnej 0...100 %RH i temperatury otoczenia od -40°C do +85°C. Zawiera interfejsy I²C i SPI.

www.bosch-sensortec.com



Seria czujników przepływu o 100 µl/min...5 ml/min do zastosowań w aparaturze medycznej

Ostatnie lata przyniosły znaczące zmiany w obszarze opieki zdrowotnej. Coraz częściej odchodzi się od hospitalizacji na rzecz leczenia w domu. Jedną z kluczowych zmian stanowi przejście z dożylnych iniekcji na podskórne za pomocą iniektorów LVI (large volume injector). Są to noszone na ciele urządzenia do ciągłego podawania leków, zmniejszające obciążenie personelu medycznego, obniżające koszty i poprawiające komfort pacjenta. Precyzyjne dawkowanie leku, kluczowe w niektórych terapiach, wymaga kontrolowania pompy za pomocą czujnika przepływu. Dodatkowo – dla zwiększenia bezpieczeństwa pacjentów i zapewnienia zgodności urządzeń z obowiązującymi regulacjami – niezbędne są: wykrywanie pęcherzyków powietrza, pustych linii, blokad oraz potwierdzenie dostarczenia leku.

Aby sprostać tym wymaganiom, firma Sensirion opracowała serię cyfrowych czujników przepływu cieczy SLD3x, łączących precyzyjny pomiar przepływu z małymi gabarytami. Poza wynikami pomiaru dostarczają one danych diagnostycznych, które mogą być stosowane do wykrywania blokad i pęcherzyków powietrza w linii, sygnalizacji awarii pompy oraz pomiaru temperatury cieczy.

Seria SLD3x obejmuje modele o zakresach pomiarowych od 100 µl/min do 5 ml/min. Charakteryzują się one dokładnością 5%, powtarzalnością 0,5% i maksymalnym ciśnieniem roboczym 3 b. Zawierają interfejs I²C do komunikacji z mikrokontrolerem. Są produkowane w obudowach o wymiarach 14×12×3,2 mm z portami procesowymi umieszczonymi na dolnej powierzchni.

www.sensirion.com



automatycznie wykrywa brak podłączenia USB i redukuje pobór prądu do 0,7 mA, a w stanie disabled – do 13 µA. Wykrywa również i raportuje podłączenie układu współpracującego oraz zakończenie sukcesem negocjacji w trybie dużej prędkości.

PI5USB212 pracuje z napięciem zasilania od 2,3 do 5,5 V. Jest produkowany w obudowie X2-QFN1616-12 o powierzchni 1,6×1,6 mm. Zapewnia odporność na wyładowania ESD do 2 kV (HBM) i 1 kV (CDM).

www.diodes.com



Miniaturowe moduły komunikacyjne LTE Cat 1bis

Firma u-blox rozszerza rodzinę modułów komunikacyjnych R10 na szybko rozwijający się rynek łączności komórkowej LTE Cat 1bis. Do oferty wchodzi dwa nowe modele, ułatwiające użytkownikom migrację ze starszej technologii komórkowej 2G/3G. LEXI-R10 Global, produkowany w obudowie o rozmiarach zaledwie 16×16 mm, nadaje się idealnie do zastosowań w urządzeniach przenośnych oraz systemów lokalizowania ludzi i zwierząt. Jest obecnie najmniejszym na rynku modulem LTE Cat 1bis z funkcją lokalizacji wewnątrz pomieszczeń. Z kolei nowe moduły SARA-R10 oferują te same parametry, co LEXI-R10 Global, występują natomiast w module popularnego formatu SARA. Oferują projektantom możliwość łatwego przejścia z technologii 2G/3G na standard 4G LTE o największym obecnie pokryciu globalnym.

LEXI-R10 i SARA-R10 mogą korzystać z modułów eSIM oraz zawierają zintegrowany Wi-Fi Sniffer, umożliwiający lokalizację wewnątrz budynków w oparciu na serwisie u-blox CellLocate, działającym zarówno w sieciach Wi-Fi, jak i komórkowych. Jeden z dostępnych wariantów, SARA-R10M10, jest najmniejszym obecnie modulem LTE Cat 1bis z wbudowanym odbiornikiem GNSS, nadającym się idealnie do aplikacji telematyki i śledzenia zasobów, wymagających ciągłej łączności. Może być stosowany w dowolnych systemach pozycjonowania, działających wewnątrz i na zewnątrz pomieszczeń, w dowolnej lokalizacji geograficznej. Jego powierzchnia jest mniejsza o połowę od wcześniejszego odpowiednika LTE Cat 1bis combo o oznaczeniu LENA-R8M10.

www.u-blox.com



Samoadaptujący się kondycjoner sygnału do portów USB 2.0

PI5USB212 to kondycjoner sygnałów USB 2.0 realizujący funkcje wzmocnienia i preemfazy w celu skompensowania strat sygnału ISI w kanałach. Może znaleźć zastosowanie w laptopach, stacjach dokujących, ekstenderach, monitorach i odbiornikach TV. Jest zgodny ze specyfikacjami USB 2.0, OTG 2.0 oraz BC 1.2.

Układ pozwala wydłużyć zasięg kabli USB 2.0 do 5 m i może być umieszczany daleko od interfejsu USB PHY lub złącza. Poziomy wzmocnienia i preemfazy oraz czułość odbiornika są konfigurowane za pomocą wyprowadzeń lub przez interfejs I²C. PI5USB212

REKLAMA

BORNICO
to miejsce, które
łącząc doświadczenie
z innowacyjnością sprawia, że Twoje
pomysły nabierają życia.

✉ bornico@bornico.com.pl
🌐 www.bornico.com.pl

☎ +48 517 312 709 | +48 517 312 419

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.



Płytką stykową Genius – kompletny warsztat elektroniczny w jednym zestawie

Breadboard Genius stanowi nowe podejście do płytek stykowych, używanych do prototypowania urządzeń elektronicznych. System ten integruje w sobie nie tylko samą płytkę do montażu elementów w czasie prototypowania, ale również zasilacz i wbudowany oscyloskop.

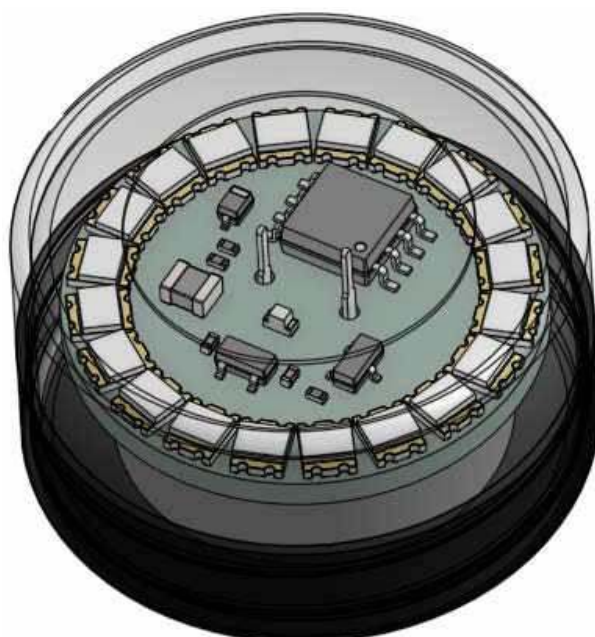
Projekt można nazwać kieszonkowym laboratorium elektronika – idealnym zarówno dla nauczycieli, jak i hobbystów. Dzięki wymiennym płytkom stykowym, umożliwia on jednoczesną pracę nad kilkoma projektami, a modułowa konstrukcja sprzyja łatwej rozbudowie. Zestaw oferuje bramki logiczne, wskaźniki LED oraz przełączniki, które pozwalają na przeprowadzanie eksperymentów czy też naukę podstaw elektroniki i testowanie bardziej zaawansowanych projektów.

Podstawowym elementem systemu jest zasilacz z wyjściami stałymi i zmiennymi od -5 V do 30 V. Drugi kluczowy element miniaturowego „laboratorium” stanowi zintegrowany oscyloskop jednokanałowy. Dwa 7-segmentowe wyświetlacze LED umożliwiają łatwe monitorowanie wartości i wyników działania wielu rodzajów układów, a całość uzupełniają 16 diod LED w różnych kolorach.

Pozostałe moduły Breadboard Genius zawierają regulowane rezystory o wartościach 1 kΩ, 10 kΩ i 100 kΩ, zintegrowane bramki

logiczne (w zestawie znajdują się bramki AND, OR, NOT, NAND, NOR i Exclusive OR), a także elementy elektromechaniczne do sterowania układem: 8 przycisków i 8 przełączników SPDT. Zasilanie całego systemu dostarcza port.

<https://tiny.pl/5np4f5n0>



Wieczny lampion fotowoltaiczny

Opisywany projekt to małe urządzenie zasilane energią słoneczną, które zawsze jest w stanie zapewnić odrobinę światła. Lampion wykonano z małego pudełka po kosmetykach – zainstalowano w nim diodę LED włączającą się w nocy (zawsze, nawet jeśli lampion trzymany jest w pomieszczeniu). Jedyne wymaganie? Dostęp do niewielkiej ilości naturalnego światła w ciągu dnia.

W opisie projektu autor nie tylko pokazuje sam lampion, ale także dokumentuje, czego się nauczył i jakie decyzje podjął podczas projektowania oraz budowy tego urządzenia.

Projekt oparty jest na mikrokontrolerze ATtiny85, który steruje prostą przetwornicą. Wprawdzie ATtiny85 nie należy do mikrokontrolerów o niskim poborze mocy, ale przez większość czasu pozostaje w stanie uśpienia, zużywając wówczas około 4 μA. Pod względem zużycia energii nie jest zatem idealnie, ale to dobry początek.

Do zbierania energii świetlnej zastosowano fotodiody. Dzięki temu cały lampion utrzymał bardzo małe wymiary – mieści się w typowym słoiczku do kosmetyków 3 g (jego średnica to niecałe 30 mm, a wysokość 16 mm). Emituje światło dostatecznie jasne, by można było je łatwo dostrzec w nocy. W całkowitej ciemności lampion świecić może nawet kilka dni.

<https://hackaday.io/project/196267-lampion>

Ulubiony Kiosk

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
multimediów
samodzielnie.**

Szczegóły na UlubionyKiosk.pl/media



W ofercie AVT*

AVT6060

Najważniejsze parametry:

- odbiór stereofoniczny audycji nadawanych analogowo w paśmie CCIR (87,5...108 MHz) z modulacją częstotliwościową (FM),
- wbudowany wzmacniacz głośnikowy 2x1 W (8 Ω, 9 V) z regulacją głośności za pomocą potencjometru,
- wyjście liniowe audio,
- prosty wskaźnik siły sygnału,
- zapamiętywanie jednej stacji w nieulotnej pamięci EEPROM,
- wyszukiwanie stacji zarówno w dół skali, jak i w górę,
- pobór prądu 20...500 mA,
- zasilanie napięciem stałym 9 V (lub 5 V – po niewielkiej modyfikacji).

* **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] – 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

- AVT5810 Miniaturowy odbiornik FM (EP 10/2020)
- Projekt 242 Moduł tunera radiowego FM (EP 4/2019)
- AVT5627 radioBox – miniaturowy odbiornik radiowy dla aktywnych (EP 5/2018)
- AVT5540 Radioodbiornik dla każdego (EP 5/2016)
- Radio – radioodbiornik stereofoniczny z RDS-em (EP 08/2015)
- AVT5401 PocketRadio – radioodbiornik kieszonekowy z RDS (EP 6-7/2013)
- AVT5317 Lampowo-tranzystorowy odbiornik UKF (EP 11/2011)
- AVT5242 Radioodbiornik internetowy (EP 7/2010)
- AVT5016 Amplituner FM z RDS (EP 6-7/2001)
- AVT2469 Odbiornik UKF FM (EdW 1/2001)
- AVT2330 Miniaturowy odbiornik FM stereo (EdW 2/1999)

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.



Bardzo prosty odbiornik UKF FM

Proste radyjka na zakres UKF, które radioamatorzy budowali jeszcze kilkanaście lat temu, wymagały specjalistycznych przyrządów do zestrojenia. Nawet te bardzo proste odbiorniki, na przykład superreakcyjne, wymagały nawijania cewek. Teraz każdy może sam zbudować swój prosty odbiornik UKF FM, który jednocześnie będzie bardzo praktyczny i intuicyjny w obsłudze.

Kilkadziesiąt lat temu zbudowanie własnego odbiornika radiowego wymagało przebrnięcia przez instrukcję, która zawierała określenia w stylu: L1 = 270 zwojów DNE 0,18 na karkasie o średnicy 6 mm. Im wyższa była częstotliwość pracy urządzenia, tym (pozornie!) prostsze wydawało się zadanie, gdyż malała liczba zwojów i drut

nawojowy stawał się grubszy. Tak naprawdę jednak utrzymanie tej cewki w ryzach wcale nie było proste. Drut się gwał, a zwoje były nierówno rozłożone, co finalnie przekładało się na problemy z zestrojeniem. Generalnie – trudny temat.

Owszem, istniała wtedy (i istnieje po dziś dzień) spora rzesza elektroników, którym

nawijanie cewek nie przeszkadza, ba, nawet traktują to jako swoisty smaczek – coś innego niż zwyczajne lutowanie. Panie i panowie – macie mój szacunek, naprawdę! Zaś tych, którzy chcą sami posłuchać, co w eterze piszczy, ale „kręcenie” własnych cewek po prostu nie kręci, zapraszam do lektury opisu niniejszego urządzenia.

Budowa układu

Schemat ideowy omawianego urządzenia znajduje się na **rysunku 1**. Jako pierwszy blok można z pewnością wyróżnić na nim moduł z układem TEA5767. Ta niewielka płytka wykonuje za nas całą czarną robotę, a co jeszcze lepsze – steruje się nią w pełni

Wykaz elementów:**Rezystory:** (THT o mocy 0,25 W)

R1, R2: 10 kΩ
R3, R4: 4,7 Ω
R5...R7: 1 kΩ
R8, R9: 3,3 kΩ
RN1: 4 × 10 kΩ SILS
PT: 2 × 10 kΩ logarytmiczny (na panel)

Kondensatory:

C1, C8, C11...C14, C16: 100 nF (raster 5 mm, MKT)

C2...C6, C15: 100 μF 16 V (raster 2,5 mm)
C7, C9, C10: 470 μF 25 V (raster 3,5 mm)

Półprzewodniki:

D1: 1N5819
LED1...LED3: zielona 5 mm matowa (np. LED F5 G)
U1: moduł z TEA5767
U2: TDA2822M (DIP8)
U3: 78L05 (TO92)
U4: ATtiny24A-PU (DIP14)

Pozostałe:

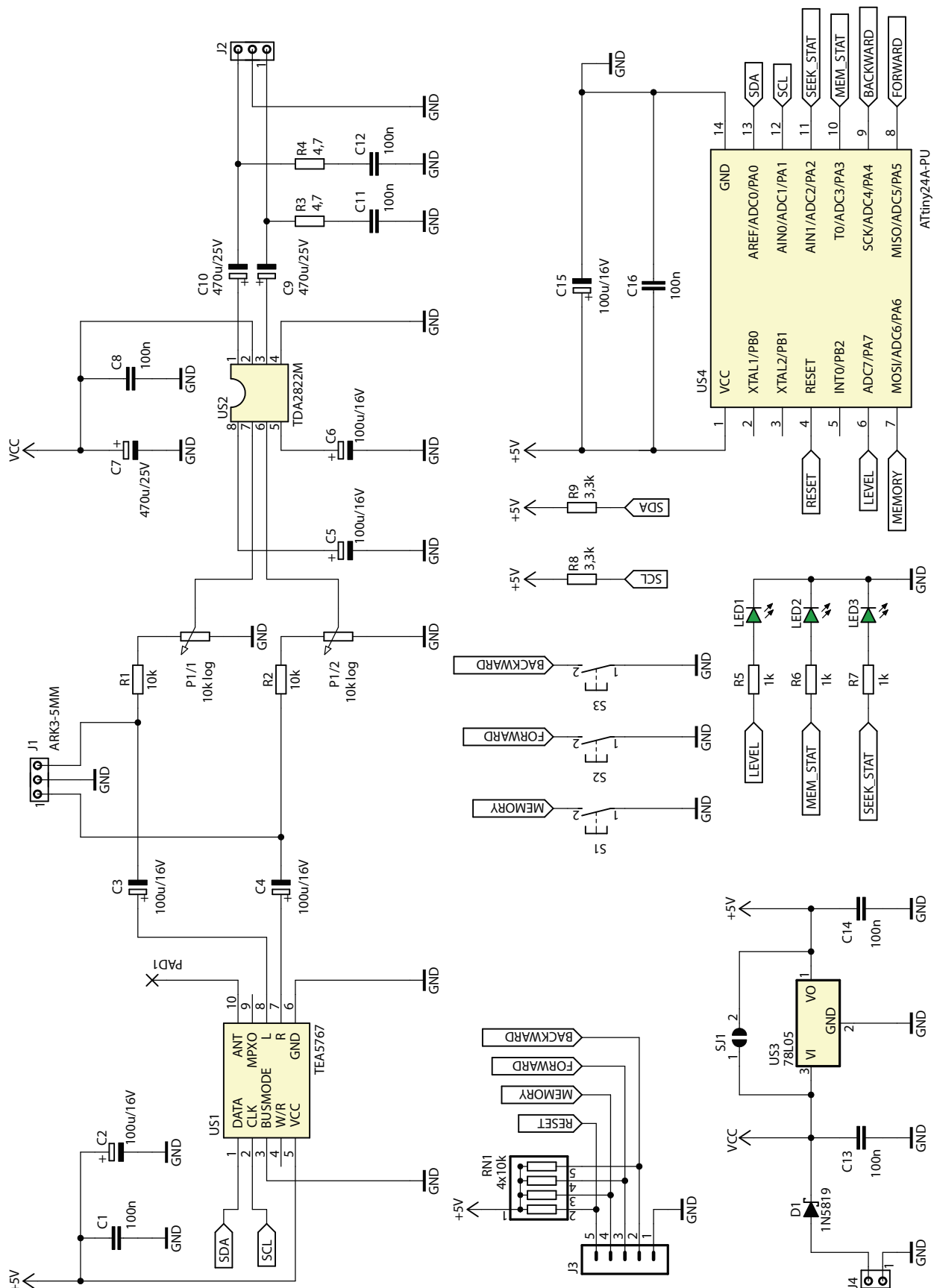
J1, J2: ARK3 (5 mm)
J3: goldpin 5 pin męski (2,54 mm, THT)
J4: ARK2 (5 mm)
S1...S3: microswitch 6×6 kątowny (MIKROSW 6 K9)
Jedna podstawka DIP8
Jedna podstawka DIP14
Antena radiowa (opis w tekście)

cyfrowo. Schemat blokowy modułu pokazuje **rysunek 2**.

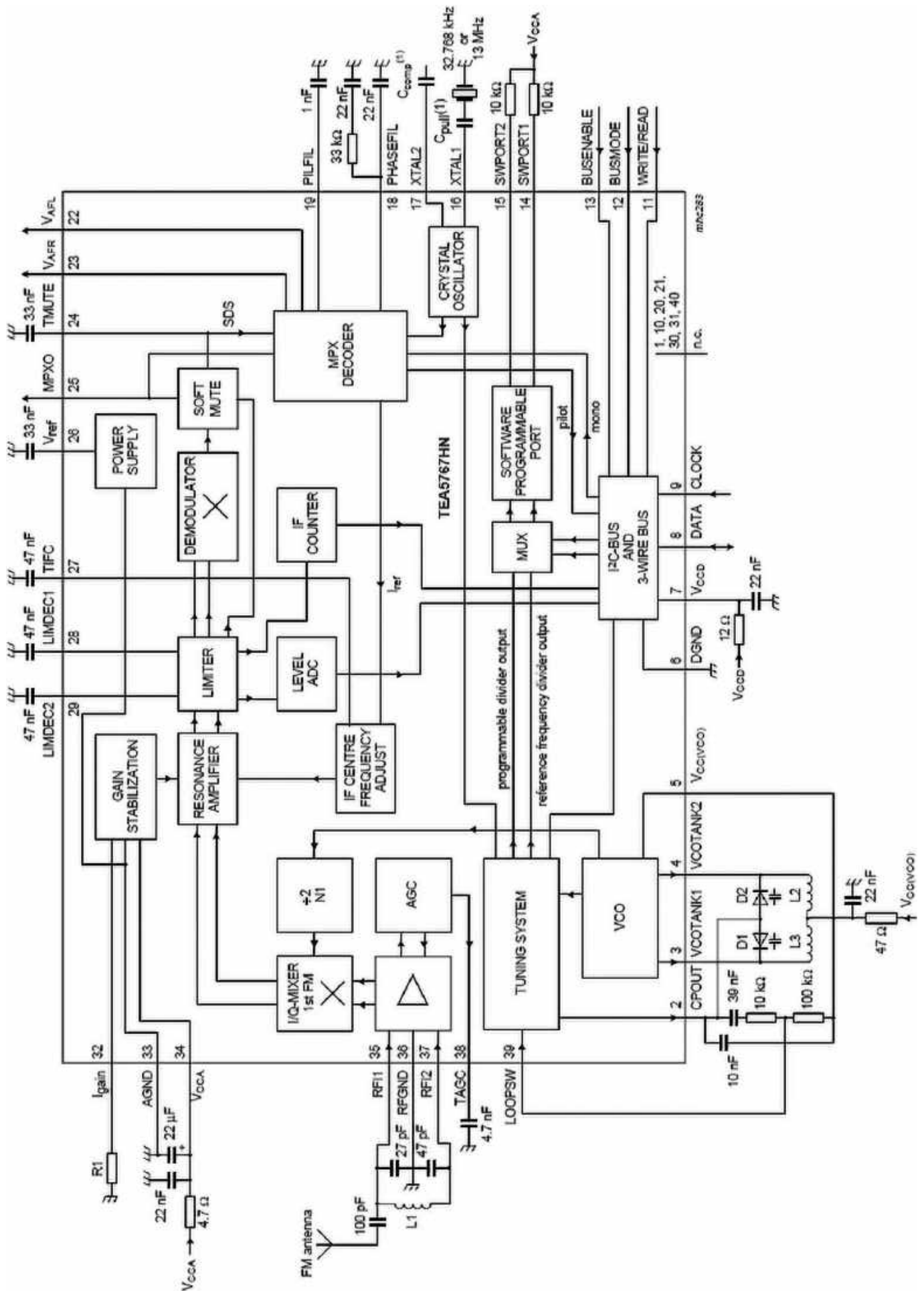
Na wejściu sygnału w.cz. (którego źródłem jest antena) znajdują się: obwód rezonansowy o niskiej dobroci (z cewką L1) oraz szerokopasmowy wzmacniacz wielkiej

częstotliwości, którego wzmocnienie jest regulowane automatycznie (AGC). Cyfrowy system strojenia, będący niczym innym, jak generatorem na bazie pętli synchronizacji fazowej (PLL), wytwarza sygnał heterodyny, korzystając z oscylatora przestrajanego na

pięciowo (VCO). Częstotliwość wzorcowa dla tego generatora jest ustalona przez dołączony z zewnątrz rezonator kwarcowy o częstotliwości 32768 Hz lub 13 MHz – w przypadku omawianego modułu prawdziwa jest pierwsza z dwóch wymienionych wartości



Rysunek 1. Schemat ideowy bardzo prostego odbiornika UKF FM



Rysunek 2. Schemat blokowy układu TEA5767

(zastosowano niewielki, cylindryczny kwarc zegarkowy).

Wzmocniony sygnał wielkiej częstotliwości oraz wytworzony przez VCO sygnał heterodyny trafiają do mieszacza kwadraturowego (I/Q). Obie powstałe w nim składowe (synfazowa I oraz kwadraturowa Q) są wzmacniane przez wzmacniacz rezonansowy, który jest niczym innym, jak znany z typowego układu superheterodyny wzmacniaczem pośredniej częstotliwości. Z tym ostatnim z kolei współpracuje ogranicznik sygnału, eliminujący niepożądaną modulację amplitudy. Przy nim znajduje się również przetwornik A/C mierzący aktualną siłę sygnału radiowego poprzez określenie jego średniej amplitudy.

Sygnał, po opuszczeniu ogranicznika amplitudy, poddawany jest demodulacji, a następnie trafia na dekodery stereo. W zależności od warunków może on pracować lub nie, ponieważ obiór stereofoniczny wymaga sygnału o lepszej jakości. Całość działa zatem jak klasyczny odbiornik superheterodynowy, z tą różnicą, że zdecydowana większość jego układu została zintegrowana na jednym płasku krzemu.

Wszystkie znajdujące się na prezentowanym schemacie cewki zostały przygotowane, przylutowane i zestrojone przez producenta modułu. Nam pozostaje jedynie przylutować dziesięć pól lutowniczych, które są wyprowadzone na dwóch bocznych krawędziach tegoż modułiku. Pole lutownicze PAD1 przewidziane jest do podłączenia anteny – o czym dalej.

Otoczenie modułu okazuje się bardzo proste: wymaga on dostarczenia zasilania DC oraz cyfrowego sygnału sterującego. Ów sygnał to pięć bajtów przesyłanych magistralą I²C. Rezystory R8 i R9 podciągają linie sygnałowe do dodatniego potencjału zasilającego (około +5 V). Na wyjściach zdemodulowanego sygnału małej częstotliwości istnieje składowa stała, toteż kondensatory C3 i C4 odcinają ją. Voilà, oto cała magia czułego odbiornika radiowego sprowadzona do kilku prostych podzespołów.

Wbudowany w płytke wzmacniacz audio został zrealizowany przy użyciu starego już układu TDA2822M, którego aplikacja jest bardzo prosta, zaś parametry – wystarczające do zastosowania go w odbiorniku

określanym jako „prosty”. Amplituda sygnału wychodzącego z TEA5767 pozostaje na tyle wysoka, że bez problemu przestawia wejście tego stereofonicznego wzmacniacza mocy. Rezystory R1 i R2 – w połączeniu z rezystancją ścieżek oporowych potencjometru P1 – tworzą dzielnik napięciowy, tłumiący poziom sygnału wejściowego o około 6 dB.

Kondensatory C5 i C6 są wymagane przez notę katalogową układu TDA2822M do jego poprawnej pracy. Ich rola polega na zwieraniu do masy wejść odwracających, lecz tylko dla sygnału – stąd użycie kondensatorów, a nie po prostu podłączenie do masy. Na wyjściach znajduje się składowa stała, którą separują od głośników kondensatory C9 i C10. Szeregowe obwody RC (C11+R3 i C12+R4) stanowią tzw. obwody Zobla, zaś ich zadaniem jest obciążanie stopni wyjściowych dla składowych o wysokiej częstotliwości – czyli tych, przy których reaktancja indukcyjna cewki głośnika staje się już bardzo wysoka.

Jeżeli ktoś chce użyć wzmacniacza audio innego niż wbudowany, może pobrać zdemodulowany sygnał analogowy z zaciski złącza J1. Wtedy jednak nie będzie dostępna regulacja głośności za pomocą znajdującego się na płytce potencjometru P1.

Pracą modułu TEA5767 steruje prosty i tani mikrokontroler typu ATtiny24-PU. Ma niewiele pamięci programu, lecz okazuje się ona wystarczająca do realizacji tego projektu. Procesor można programować w układzie, bez wyciągania z podstawki, poprzez złącze J3 (to na nie zostały wyprowadzone linie sygnałowe interfejsu ISP). Linie przeznaczone do programowania zostały dodatkowo podciągnięte do potencjału +5 V, aby zlikwidować problem gromadzenia się na nich ładunków elektrostatycznych oraz aby zmniejszyć ich impedancję dla zakłóceń elektromagnetycznych. Dodatkowo owe trzy linie zostały zaprzęgnięte do sprawdzania stanu przycisków monostabilnych S1...S3.

Użytkownik, poza przyciskami do sterowania odbiornikiem radiowym, ma do dyspozycji również trzy diody LED sterowane przez mikrokontroler. Rezystory R5...R7 ograniczają ich prąd przewodzenia do około 3 mA, co stanowi wartość wystarczającą do wyraźnego świecenia, zwłaszcza w przy-

padku diod zielonych (na tę barwę ludzkie oko jest najbardziej wyczulone). Zasilanie do układu doprowadzają zaciski złącza J4. Dioda D1 chroni podzespoły przed uszkodzeniem w razie odwrotnego podłączenia zewnętrznego zasilacza. Stabilizator 78L05 dostarcza napięcia 5 V do modułu radiowego oraz mikrokontrolera, zaś wzmacniacz mocy jest zasilany wprost z katody D1.

Montaż i uruchomienie

Układ został zmontowany na dwustronnej płytce drukowanej o wymiarach 110 mm × 35 mm. Wzór jej ścieżek oraz schemat montażowy pokazuje **rysunek 3**. Wszystkie cztery otwory montażowe mają średnicę 3,2 mm i zostały umieszczone w odległości 3 mm od krawędzi płytki.

Montaż proponuję rozpocząć od elementów o najmniejszej wysokości obudowy, czyli modułu radiowego US1 oraz rezystorów. Pod układy US2 i US4 proponuję zastosować podstawki, aby ułatwić wymianę któregoś z nich (w razie uszkodzenia) – oraz programowanie mikrokontrolera. W układzie prototypowym diody LED1...LED3 zostały wlutowane na zgiętych nóżkach, aby mogły wystawać przez przednią ściankę obudowy – tak jak klawisze przycisków S1...S3 i oś potencjometru P1. Zmontowany układ można zobaczyć na **fotografii tytułowej**.

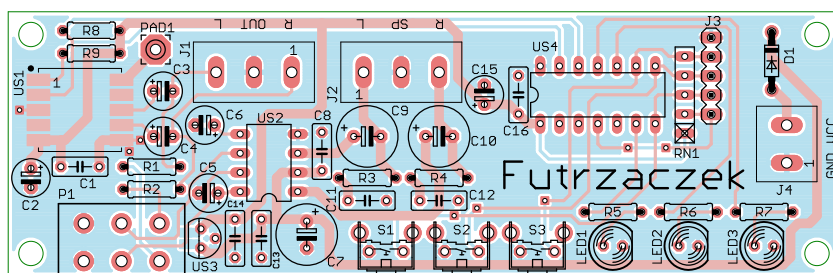
Na etapie uruchamiania konieczne jest zaprogramowanie pamięci Flash mikrokontrolera dostarczonym wsadem oraz zmiana jego bitów zabezpieczających. Oto ich nowe wartości:

Low Fuse = 0xE2

High Fuse = 0xDC

Szczegóły są widoczne na **rysunku 4**, który zawiera widok okna konfiguracji tychże bitów w programie BitBurner. W ten sposób zostanie uruchomiony wewnętrzny generator RC o częstotliwości 8 MHz (a dokładniej: wyłączony prescaler przez 8) oraz Brown-Out Detector, który wprowadzi mikrokontroler w stan zerowania, jeżeli jego napięcie zasilające spadnie poniżej 4,3 V. Taki zabieg znacznie zmniejsza ryzyko zawieszenia się mikrokontrolera podczas uruchamiania.

Poprawnie zaprogramowany układ jest gotowy do działania po podłączeniu zasilania



Rysunek 3. Schemat płytki PCB



Rysunek 4. Szczegóły ustawienia bitów zabezpieczających

do zacisków złącza J4. Powinno to być napięcie stałe, dobrze filtrowane, najlepiej stabilizowane – na przykład z zasilacza wtyczkowego. W podstawowej wersji układu, z wlutowanym stabilizatorem US3, powinno ono wynosić 9 V. Górny limit tego napięcia, wyznaczony przez układ TDA2822M, sięga 15 V, lecz przy tak wysokim napięciu zasilającym i pracy z niskoomowymi głośnikami dojdzie do przegrzania wzmacniacza. W tych warunkach rekomenduję użycie głośników o impedancji 8 Ω , a najlepiej jeszcze większej. Maksymalna moc wyjściowa powinna wynosić 1 W na kanał. Pobór prądu zależy od aktualnegoysterowania i może zawierać się w przedziale od 20 mA do nawet 500 mA (przy głośnikach 8 Ω).

Można też pominąć stabilizator 78L05 i cały układ zasilac napięciem stabilizowanym o wartości 5 V, na przykład z ładowarki USB. Wzmacniacz mocy znajdujący się na płytce jest przystosowany do pracy w takich warunkach, co więcej: można będzie zastosować głośniki o niższej impedancji obciążenia, na przykład 4 Ω , bez ryzyka jego przegrzania. W tym celu trzeba wylutować US3 i zewrzeć kropłą cyny dwa prostokątne pola lutownicze SJ1, które znajdują się na spodniej stronie płytki. Wówczas na cały układ trafi napięcie pochodzące wprost z zasilacza. Pobór prądu

będzie zawierać się w takich samych granicach (dla głośników 4 Ω), choć moc wyjściowa może być o około 50% niższa niż w poprzednim wariancie.

Gdyby wbudowany wzmacniacz nie był nikomu do szczęścia potrzebny, można skorzystać z innego, zewnętrznego – wtedy konieczne okaże się zastosowanie stabilizatora US3, lecz warto przy tym wyjąć z podstawki układ US2. Napięcie zasilające może wynosić 9...24 V, zaś pobór prądu będzie stały i wyniesie około 15 mA, niezależnie odysterowania. Jeżeli zaś do dyspozycji jest gotowe napięcie 5 V, wówczas można US3 pominąć (wylutować) i zewrzeć jego wejście z wyjściem zworką SJ1.

Jako anteny odbiorczej można użyć odciinka miedzianego przewodu w izolacji, który należy przylutować do pola lutowniczego PAD1, zlokalizowanego nieopodal modułu US1. Najlepiej, aby jego długość wynosiła około 1,5 m (lub połowę tej wartości, jeżeli nie ma warunków na użycie przewodu półtorametrowego). Można też użyć przeznaczonej specjalnie do zakresu UKF anteny teleskopowej.

Zaraz po włączeniu zasilania układ odczytuje częstotliwość stacji zapisanej w pamięci EEPROM. Jeżeli pamięć jest pusta lub wystąpi błąd odczytu, ustawiona zostanie najniższa możliwa częstotliwość dla

tego zakresu (87,5 MHz). Przeszukiwać pasmo można zarówno w górę (wciskając na chwilę S2), jak i w dół (po naciśnięciu S3). W trakcie wyszukiwania odpowiednio silnych stacji dioda LED3 świeci się, po zakończeniu wyszukiwania – gaśnie.

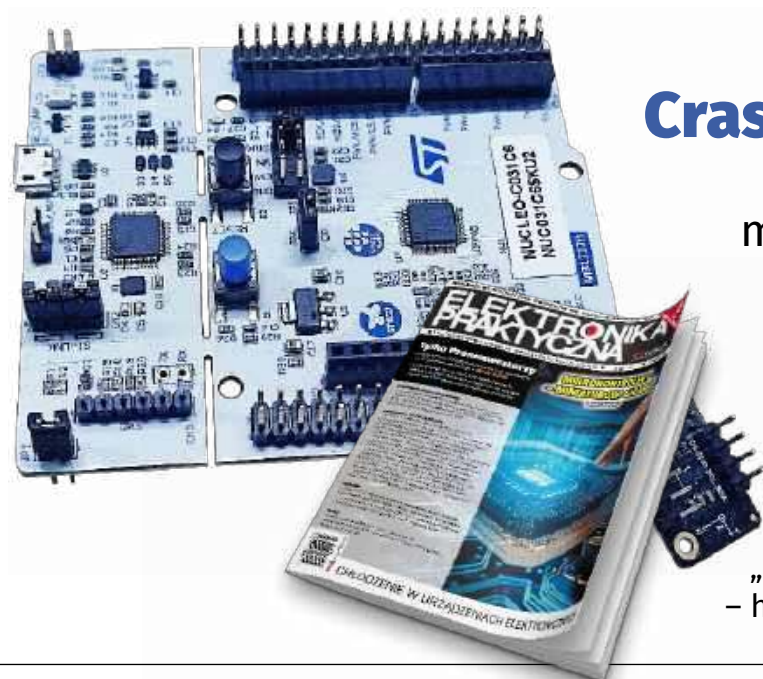
Jasność świecenia diody LED1 wskazuje na poziom odbieranego sygnału. Układ rozróżnia 16 wartości siły sygnału, zaś wypełnienie sygnału PWM sterującego diodą rośnie logarytmicznie, zatem można wyrażnie zauważyć przyrost lub spadek jasności świecenia. Ten prosty wskaźnik siły sygnału jest odświeżany po każdym wyszukiwaniu stacji oraz cyklicznie (co 5 sekund) w trakcie pracy układu.

Przycisk S1 służy do obsługi pamięci stacji. Jeżeli zostanie na chwilę wciśnięty i zwolniony, układ bezwarunkowo ustawi częstotliwość zapisanej stacji. Jeżeli zaś wciśnięcie będzie trwało dłużej niż 1 sekundę, układ zapisze tę częstotliwość do pamięci EEPROM. W sytuacji, gdy aktualnie ustawiona częstotliwość pokrywa się z zapisaną w pamięci, świeci się dioda LED2 – będzie się zatem świeciła zaraz po włączeniu zasilania, po użyciu przycisku S1 oraz kiedy wyszukiwanie stacji natrafi na tę już wcześniej zapamiętaną.

Michał Kurzela, EP

REKLAMA

Programowanie prostszych mikrokontrolerów (np. AVR, PIC, MSP430 czy też przestarzałych już 8051 bądź HCS08) bez użycia bibliotek, tj. przy wykorzystaniu samych tylko plików nagłówkowych z definicjami rejestrów i zawartych w nich bitów, jest raczej naturalną konsekwencją nieskomplikowanej architektury tych procesorów. Bardziej rozbudowane układy – w szczególności te oparte na rdzeniach ARM – są zwykle nieporównanie trudniejsze do opanowania na niskim poziomie abstrakcji, stąd większość programistów systemów wbudowanych korzysta w swojej codziennej pracy z bibliotek. Niniejszy kurs ma na celu pokazanie innej ścieżki rozwoju i – mamy nadzieję – przekona przynajmniej część spośród naszych Czytelników do zaprzyjaźnienia się z wymagającą, ale niezwykle wartościową metodą programowania układów STM32.



Crash Course STM32C0

– programowanie
mikrokontrolerów ARM
w rejestrach

Kupisz i przeczytasz
w marcowym wydaniu
„Elektroniki Praktycznej”
– <https://ulubionykiosk.pl>





W ofercie AVT*

AVT6054

Najważniejsze parametry:

- napięcie zasilania: 5 V (DC),
- maksymalny prąd obciążenia (bez podłączonych diod LED): 70 mA,
- zakres mierzonych temperatur: 0...125°C,
- rozdzielczość pomiaru temperatury: 1°C,
- dokładność pomiaru temperatury: ±0,5°C,
- zakres ustawień temperatury: 0...180°C,
- rozdzielczość ustawień temperatury: 10°C,
- zgodność ze standardami (nazw własnych użyto wyłącznie w celu identyfikacji produktu): Asus Aura Sync, Gigabyte RGB Fusion Ready, MSI Mystic Light Sync, ASRock Polychrome Sync.

* **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] – 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

- AVT6054 argbController (1) (EP 9/2024)
- AVT5784 Wolnozmienny sterownik taśmy RGB (EP 8/2020)
- AVT5778 Sterownik taśm LED RGB+W zgodny z HomeKit (EP 7/2020)
- AVT5596 Mieszacz kolorów RGB (EP 7/2017)

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.

argbController (2)

W poprzedniej części artykułu opisaliśmy szczegóły techniczne nieco tajemniczego protokołu ARGB Gen2. Tym razem przechodzimy do kwestii praktycznych – prezentujemy zarówno oprogramowanie sterujące łańcuchami LED zgodnymi z omawianym standardem transmisji, jak i konstrukcję prostego urządzenia mikroprocesorowego implementującego ów protokół.

Oprogramowanie mikrokontrolera

Plik nagłówkowy, pokazany na **listingu 1**, definiuje główne ustawienia sprzętowe, a także wprowadza niezbędne typy danych – w tym configType, który upraszcza późniejszą konfigurację diod LED. Dalej, na **listingu 2**, pokazano kod funkcji inicjalizacyjnej interfejsu komunikacyjnego, którego wyłącznym zadaniem jest ustawienie kierunku portu danych i jego stanu spoczynkowego (0).

Z kolei na **listingu 3** pokazano ciała funkcji odpowiedzialnych za przesłanie jednego bitu danych (0 lub 1), przy udziale wspomnianego wcześniej interfejsu komunikacyjnego. Warto zauważyć, że stosowne funkcje zdefiniowano, używając atrybutu `__always_inline__`, który instruuje kompilator, by ZAWSZE wstawiał w miejscu wywołania rozwinięcie funkcji, nie zaś skok do jej ciała. Jest to o tyle istotne, że operujemy na dość rygorystycznych timingach (rzędu setek nanosekund). W takiej sytuacji – przy częstotliwości taktowania mikrokontrolera równej 20 MHz (okres 50 ns) – jakiegokolwiek niepotrzebne (z punktu widzenia programisty, nie kompilatora) wykonanie rozkazu asemblera może rozłożyć całą transmisję danych, całkowicie uniemożliwiając sterowanie wspomnianymi elementami.

Skoro mamy już funkcje umożliwiające przesłanie jednego bitu informacji, pora na funkcję (także odpowiednio zadeklarowaną), umożliwiającą przesłanie kompletnego bajtu danych za pomocą interfejsu ARGB Gen2 – patrz **listing 4**.

W dalszej kolejności, by uprościć zapis danych do adresowalnych diod LED RGB, przewidziano funkcję pozwalającą na przesłanie kompletnej informacji o kolorze tejże diody LED. Ciało tej funkcji pokazano na **listingu 5**.



Pierwsza część artykułu znajduje się pod adresem: <https://ulubionykiosk.pl/media>

```
//Definicja typu strukturalnego przechowującego konfigurację diod LED
typedef struct
{
    uint8_t pwmFreq;
    uint8_t ovCurr;
    uint8_t respType;
    uint8_t fineR;
    uint8_t fineG;
    uint8_t fineB;
}configType;

//Definicje portu sterującego diodami LED
#define ARGB_PORT_NAME PORTB
#define ARGB_PORT_MASK PIN2_bm //PB2

#define ARGB_PORT_AS_OUTPUT ARGB_PORT_NAME.DIRSET = ARGB_PORT_MASK
#define ARGB_PORT_AS_INPUT ARGB_PORT_NAME.DIRCLR = ARGB_PORT_MASK
#define ARGB_PORT_SET ARGB_PORT_NAME.OUTSET = ARGB_PORT_MASK
#define ARGB_PORT_RESET ARGB_PORT_NAME.OUTCLR = ARGB_PORT_MASK
#define ARGB_PORT_IS_SET (ARGB_PORT_NAME.IN & ARGB_PORT_MASK)
#define ARGB_PORT_IS_RESET (! (ARGB_PORT_NAME.IN & ARGB_PORT_MASK))

//Definicje timingów
#define NOP __asm__ __volatile__ ("nop") //50ns
#define T0H NOP; NOP; NOP; NOP //Minimum: 200 ns
#define T0L T0H; T0H; T0H; T0H //Minimum: 800 ns
#define T1H T0H; T0H; T0H; NOP //Minimum: 650 ns
#define T1L NOP; NOP; NOP; NOP //Minimum: 200 ns
#define RESET_TIME 250 //us
#define READ_TIMEOUT 160 //us
#define COMMAND_TIMEOUT 60 //us

//Definicje stałych konfiguracji diod LED
#define PWM_FREQ_1k 0
#define PWM_FREQ_2k 1
#define PWM_FREQ_9k 2
#define PWM_FREQ_18k 3
#define OVERALL_CURR_NORMAL 0
#define OVERALL_CURR_HALVED 1
#define RESPONSE_IS_ID 1
#define RESPONSE_IS_CURRENT 0

//Definicje stałych rozkazów oraz statusów ich wykonania
#define CMD_ASSIGN_ADDRESS 0x10
#define CMD_RESET_ADDRESS 0x20
#define CMD_PING_ADDRESS 0x30
#define CMD_ACTIVATE_ADDRESS 0x40
#define CMD_EXECUTED 0
#define CMD_DISCARDED 1
```

Listing 1. Plik nagłówkowy modułu obsługi diod ARGB Gen2

Ustawienia Fuse-bitów:

FREQSEL[1:0]: 10*

RSTPINCFG[1:0]: 01*

SUT[2:0]: 111*

EESAVE: 1

*ustawienie domyślne producenta

Na razie wszystko wydaje się proste, gdyż mamy do czynienia z funkcjami niczym nieróżniącymi się od specyfikacji dobrze znanych diod WS2812. To prawda – wspominałem przecież na wstępie, że diody ARGB Gen2 są wstecznie zgodne ze specyfikacją elementów WS2812, dlatego mogą być obsługiwane przez sterowniki starego typu. Przejdźmy zatem do zagadnień nieco bardziej skomplikowanych. Na **listingu 6** pokazano ciało funkcji pozwalającej na policzenie elementów LED znajdujących się na magistrali i korzystającej z trybu odczytu protokołu ARGB Gen2. Dalej, na **listingu 7**, widzimy z kolei ciało funkcji pozwalającej na konfigurację diod LED typu ARGB Gen2.

Co ważne, aby stosowną konfigurację przeprowadzić w sposób szybki i łatwy, najlepiej użyć zmiennej typu `configType` (będącej argumentem wywołania funkcji) i predefiniowanych w pliku nagłówkowym stałych. Należy również podkreślić, że prezentowana funkcja zakłada wysłanie takiej samej konfiguracji do wszystkich diod LED na magistrali, co zwykle okaże się najbardziej pożądane. Natomiast na **listingu 8** pokazano kluczową funkcję interfejsu ARGB Gen2, stosowaną w przypadku pracy magistrali danych w topologii gwiazdy – jej zadaniem jest wysłanie rozkazu sterującego (oraz opcjonalny odbiór odpowiedzi).

Przejdźmy teraz do założeń aplikacyjnych sterownika `argbController`. Przyznam szczerze, że pomysł (choć bardziej pasuje mi tu słowo „zapotrzebowanie”) na ten projekt podsunął mi mój kolega Andrzej. Używając różnego rodzaju podzespołów PC (w tym wyposażonych w podświetlenie ARGB Gen2), zauważył on brak systemu, który w zależności od wyników pomiaru temperatury sterowałby kolorem tego rodzaju podświetlenia. Istnieją co prawda systemy

```
void argbInit(void)
{
    //Port sterujący LED, jako wyjściowy ze stanem 0
    ARGB_PORT_RESET;
    ARGB_PORT_AS_OUTPUT;
}
```

Listing 2. Kod funkcji inicjalizacyjnej interfejsu ARGB Gen2

```
//Funkcje zakładają częstotliwość taktowania mikrokontrolera równą
//F_CPU = 20 MHz (NOP = 50 ns)

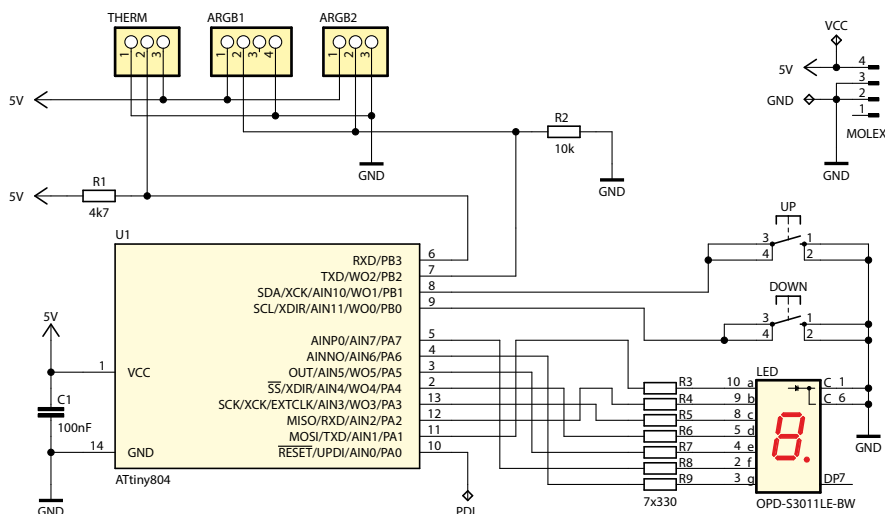
__inline__ void __attribute__((__always_inline__)) argbSendBit0(void)
{
    ARGB_PORT_SET;
    T0H;
    ARGB_PORT_RESET;
    T0L;
}

__inline__ void __attribute__((__always_inline__)) argbSendBit1(void)
{
    ARGB_PORT_SET;
    T1H;
    ARGB_PORT_RESET;
    T1L;
}
```

Listing 3. Funkcje odpowiedzialne za przesłanie jednego bitu danych interfejsu ARGB Gen2

```
__inline__ void __attribute__((__always_inline__)) argbSendByte(uint8_t Byte)
{
    ATOMIC_BLOCK(ATOMIC_RESTORESTATE)
    {
        if(Byte & 0x80) argbSendBit1(); else argbSendBit0();
        if(Byte & 0x40) argbSendBit1(); else argbSendBit0();
        if(Byte & 0x20) argbSendBit1(); else argbSendBit0();
        if(Byte & 0x10) argbSendBit1(); else argbSendBit0();
        if(Byte & 0x08) argbSendBit1(); else argbSendBit0();
        if(Byte & 0x04) argbSendBit1(); else argbSendBit0();
        if(Byte & 0x02) argbSendBit1(); else argbSendBit0();
        if(Byte & 0x01) argbSendBit1(); else argbSendBit0();
    }
}
```

Listing 4. Funkcja odpowiedzialna za przesłanie jednego bajtu danych interfejsu ARGB Gen2



Rysunek 6. Schemat ideowy urządzenia `argbController`

```
__inline__ void __attribute__((__always_inline__)) argbSendColor(uint8_t R, uint8_t G, uint8_t B)
{
    argbSendByte(G);
    argbSendByte(R);
    argbSendByte(B);
}
```

Listing 5. Funkcja pozwalająca na przesłanie kompletnej informacji o kolorze diody LED typu ARGB Gen2

Wykaz elementów:

Rezystory: (SMD 0805)

R1: 4,7 kΩ

R2: 10 kΩ

R3...R9: 330 Ω

Kondensatory: (SMD 0805)

C1: 100 nF (ceramiczny X7R)

Półprzewodniki:

U1: ATtiny804 (SO14)

LED: wyświetlacz 7-segmentowy 0,3" typu OPD-S3011LE-BW (wspólna katoda)

Pozostałe:

UP, DOWN: microswitch TACT, h=9 mm

ARGB2, THERM: listwa kołkowa goldpin, prosta, 1×3 pin

ARGB1: listwa kołkowa goldpin, prosta, 1×4 pin

MOLEX: gniazdo MOLEX męskie typu C5080WR-F-04P (JOINT TECH)

SONDA: termometr scalony typu DS18B20 (lub sonda w niego wyposażona)

```

uint8_t argbCountLeds(void)
{
    uint8_t Timeout = 0;
    uint8_t Leds = 0;

    ATOMIC_BLOCK(ATOMIC_RESTORESTATE)
    {
        //Inicjujemy tryb odczytu
        ARGB_PORT_SET;
        _delay_us(20);
        ARGB_PORT_RESET;
        _delay_us(10);
        ARGB_PORT_SET;
        _delay_us(50);
        ARGB_PORT_RESET;
        //Konfigurujemy interfejs danych, jako wejściowy, by odczytywać impulsy wysyłane przez diody LED
        ARGB_PORT_AS_INPUT;

        //Odczytujemy kolejne impulsy odpowiedzi wysyłane przez diody LED. Magistrala podciągnięta jest do masy rezystorem 10k
        while(1)
        {
            //Sprawdzamy czy nie wystąpił time-out oznaczający koniec transmisji lub zupełny brak diod na magistrali
            if(ARGB_PORT_IS_RESET)
            {
                if(++Timeout > READ_TIMEOUT) break;
                _delay_us(1);
            }
            //Liczymy impulsy wysyłane przez diody LED by określić liczbę tychże diod na magistrali danych
            if(ARGB_PORT_IS_SET)
            {
                ++Leds;
                //Czekamy na zakończenie bieżącego impulsu (stanu wysokiego)
                while(ARGB_PORT_IS_SET);
                _delay_us(1);
                Timeout = 0;
            }
        }
        //Konfigurujemy interfejs danych, jako wyjściowy (domyślnie)
        ARGB_PORT_AS_OUTPUT;
    }

    return Leds;
}

```

Listing 6. Funkcja pozwalająca na policzenie elementów LED znajdujących się na magistrali ARGB Gen2

```

void argbSendConfig(configType *Config, uint8_t Leds)
{
    uint8_t Byte1, Byte2, Byte3;

    //Przygotowujemy dane do wysłania według specyfikacji słów konfiguracyjnych
    Byte1 = (Config->pwmFreq<<6)|(Config->fineG;
    Byte2 = (Config->ovCurr<<7)|(Config->ovCurr<<6)|(Config->respType<<5)|(Config->fineR;
    Byte3 = (Config->ovCurr<<7)|(Config->fineB;

    ATOMIC_BLOCK(ATOMIC_RESTORESTATE)
    {
        //Inicjujemy tryb konfiguracji
        ARGB_PORT_SET;
        _delay_us(20);
        ARGB_PORT_RESET;
        _delay_us(300);

        //Wysyłamy taką samą konfigurację do wszystkich diod LED
        for(uint8_t i=0; i<Leds; i++) argbSendColor(Byte2, Byte1, Byte3);
        _delay_us(300);
    }
}

```

Listing 7. Funkcja pozwalająca na konfigurację diod LED typu ARGB Gen2

pozwalające zarządzać bogactwem różnych efektów RGB (np. Asus Aura Sync, Gigabyte RGB-Fusion Ready, MSI Mystic Light Sync czy ASRock Polychrome Sync), lecz nie oferują one tak pożądaną, a jednocześnie prostą funkcjonalności, jak zmiana koloru podświetlenia na skutek zmian temperatury elementu monitorowanego. Łatwo wyobrazić sobie zastosowanie takiego systemu do monitorowania stanu radiatora procesora PC lub temperatury wewnątrz obudowy typu desktop: w ramach pełnionej funkcji zmieniałby on (w sposób sugestywny) kolor podświetlenia wentylatora, dając bezpośrednią informację na temat kondycji urządzenia. To oczywiście tylko jedno z wielu możliwych zastosowań, gdyż nic nie każe nam ograniczać się do sprzętu z branży PC.

Schemat ideowy układu pokazano na **rysunku 6**. Sercem urządzenia jest

mikrokontroler ATtiny804, taktowany wewnętrznym oscylatorem o częstotliwości

równej 20 MHz i odpowiedzialny za realizację całej założonej funkcjonalności urządzenia. Jako element pomiarowy (termometr) zastosowano scalony, cyfrowy i bardzo dobrze znany element typu DS18B20 produkcji Maxim Integrated, którego obsługę zrealizowano przy użyciu programowej implementacji magistrali 1-Wire. Wybór tego elementu z szerokiej palety dostępnych termometrów scalonych podyktowany był jego dużą dostępnością i niską ceną. Co więcej: w handlu znaleźć możemy gotowe moduły w formie wodoszczelnych sond pomiarowych, odpornych na warunki zewnętrzne, z zatopionym wewnątrz układem DS18B20, przez co wybór ten stał się jeszcze bardziej oczywisty. Ponadto mikrokontroler nasz obsługuje jednocyfrowy,

REKLAMA

Hurtownia elementów elektronicznych "AKSOTRONIK" zaprasza do swojego sklepu internetowego. Zaloguj się i kupuj OEN-LJSE na naszej stronie: WWW.AKSOTRONIK.COM.PL

Aksotronik
ELEMENTY ELEKTRONICZNE

Koszulki elektroniczne zaciskowe Ceny od 0.22zł

Magnesy neodymowe oraz ferrytowe Ceny od 0.10zł

Przełączniki klawiszowe wodoszczelne i polioszczelne Ceny od 2.40zł

Druki opornikowe od 0.16 do 0.33m Ceny od 5.70zł

Przewodniki do przewodów Ceny od 11.00zł

Szczotki węglowe do elektronarzędzi Ceny od 2.40zł

Przełączniki do elektronarzędzi zwykłe i elektromagnetyczne Ceny od 7.00zł

Złącza hermetyczne Superseal Ceny od 1.10zł

Pudełka organizery Ceny od 0.95zł

Zestawy śrubek M2, M3 z nakrętkami i podkładkami Ceny od 1.50zł

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 destansowe, płytki, przełączniki, elementy akustyczne, rezystory, kondensatory, ławce, podkładki, moduły Arduino. Zapraszamy do kontaktu: INFO@aksotronik.com.pl, tel: (22) 783-20-51


```

uint8_t argbSendCommand(uint8_t Command, uint8_t Address)
{
    uint8_t cmdStatus = CMD_DISCARDED, Timeout = 0;

    ATOMIC_BLOCK(ATOMIC_RESTORESTATE)
    {
        //Inicjujemy tryb rozkazów
        ARGB_PORT_SET;
        _delay_us(40);
        ARGB_PORT_RESET;
        _delay_us(12);

        //Wysyłamy rozkaz (korzystając z predefiniowanych wartości) wraz z pożądanym adresem łańcucha LED
        argbSendByte(Command|Address);

        //Sprawdzamy odpowiedź łańcucha, która to powinna być wysłana przez pierwszą z diod LED w tymże łańcuchu LED (tak zwaną brakmę)
        //Odpowiedź ta to impuls stanu wysokiego o długości 60uS wysłany w losowym czasie od 10uS do 60uS od zbocza opadającego ostatniego
        //bitu rozkazu. Konfigurujemy interfejs danych, jako wejściowy, by odczytać odpowiedź łańcucha biorąc pod uwagę ewentualny time-out
        ARGB_PORT_AS_INPUT;

        while(1)
        {
            //Sprawdzamy czy nie wystąpił time-out oznaczający niewykonanie rozkazu
            if(ARGB_PORT_IS_RESET)
            {
                if(++Timeout > COMMAND_TIMEOUT) break;
                _delay_us(1);
            }
            //Sprawdzamy obecność impulsu stanu wysokiego wysłanego przez bramkę łańcucha LED potwierdzającego wykonanie rozkazu
            if(ARGB_PORT_IS_SET)
            {
                //Czekamy na zakończenie bieżącego impulsu (stanu wysokiego)
                while(ARGB_PORT_IS_SET);
                _delay_us(1);
                cmdStatus = CMD_EXECUTED;
                break;
            }
        }
        //Konfigurujemy interfejs danych, jako wyjściowy (domyślnie)
        ARGB_PORT_AS_OUTPUT;
    }

    return cmdStatus;
}

```

Listing 8. Funkcja pozwalająca na przestanie rozkazu i opcjonalny odbiór odpowiedzi na magistrali typu ARGB Gen2

7-segmentowy wyświetlacz LED oraz dwa przyciski funkcyjne, oznaczone jako UP i DOWN, które stanowią elementarną wersję interfejsu użytkownika. W konstrukcji układu przewidziano również złącza goldpin przeznaczone do podłączenia wentylatorów (lub innych urządzeń), wyposażonych w interfejs ARGB Gen2 zasilany napięciem 5 V. Zastosowano 2 typy złączy w różnej konfiguracji sygnałów wyjściowych, typowych dla płyt głównych ASUS/ASROCK/MSI lub GIGABYTE. Do zasilania sterownika napięciem 5 V przewidziano z kolei typowe gniazdo MOLEX, stosowane w komputerach klasy PC, gdyż kable tego rodzaju – wyposażone we wtyki żeńskie – znajdują się w każdym komputerze typu desktop (i służą m.in. do zasilania napędów). Do obsługi przycisków funkcyjnych oraz programowej eliminacji drgania styków zastosowano przerwanie od przepelnienia timera TCBO wbudowanego w strukturę mikrokontrolera, więc możliwa stała się również detekcja krótkiego i długiego naciśnięcia przycisków – bez blokowania programu obsługi aplikacji.

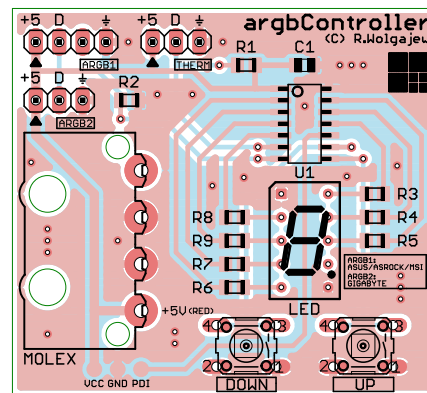
Montaż i uruchomienie

Schemat montażowy naszego urządzenia pokazano na **rysunku 7**. Niewielki, dwustronny obwód drukowany przeznaczony jest w zdecydowanej większości do montażu powierzchniowego, który rozpoczynamy od przylutowania mikrokontrolera. Element ten – mimo obudowy SMD – charakteryzuje się dość szerokim rozstawem wyprowadzeń,

przez co nie powinien narażać problemów podczas lutowania. Następnie montujemy elementy bierne, wyświetlacz LED, by na końcu wlutować złącza goldpin ARGB1, ARGB2 i THERM, przyciski UP i DOWN oraz złącze zasilające MOLEX. W przypadku złącza goldpin ARGB1 należy usunąć nieużywany pin drugi od prawej (w celu zabezpieczenia przed odwrotnym podłączeniem), gdyż gniazda tego rodzaju nie przewidują jego obecności. Poprawnie zmontowany układ nie wymaga jakichkolwiek regulacji i powinien działać tuż po włączeniu zasilania i zaprogramowaniu procesora. Na **fotografii tytułowej** zaprezentowano wygląd zmontowanej płytki drukowanej urządzenia, widzianej od strony TOP.

Obsługa urządzenia

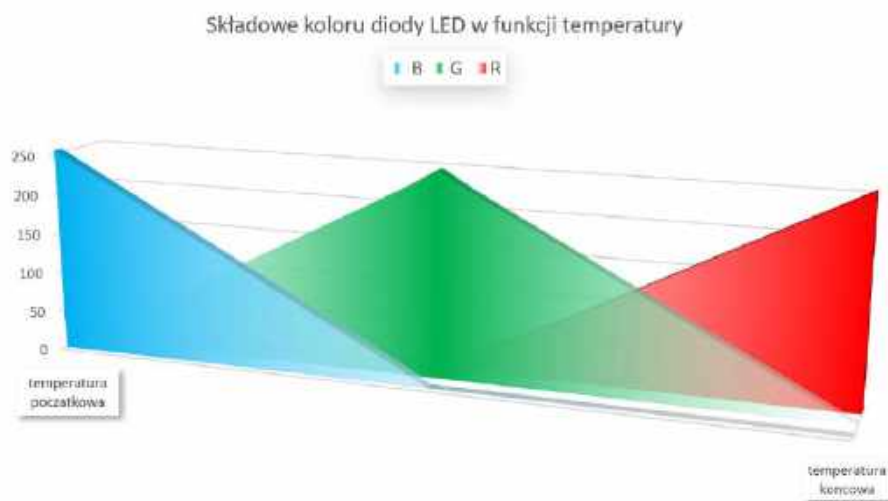
Obsługa sterownika argbController jest niezmiernie prosta. Do ustawienia mamy dwie wartości: temperaturę początkową i przedział temperatur, przy czym obie wartości regulowane są ze skokiem 10°C, w zakresie od 0°C do 90°C dla temperatury początkowej i od 10°C do 90°C – dla przedziału temperatur. Co ważne, na wyświetlaczu urządzenia temperatury pokazywane są w dziesiątkach °C. Na podstawie wprowadzonych ustawień urządzenie wylicza składowe kolorów RGB i wysyła je (takie same) do wszystkich diod LED we wszystkich podłączonych łańcuchach LED, w cyklach 1-sekundowych (przy czym liczbę podłączonych łańcuchów diod LED oraz



Rysunek 7. Schemat montażowy urządzenia argbController

liczbę samych diod LED w każdym z tych łańcuchów urządzenie sprawdza wyłącznie podczas włączania zasilania). Jak łatwo się domyślić, temperatura początkowa określa minimalną temperaturę, dla której urządzenie prześle do diod LED składowe kolory RGB odpowiedzialne za wyświetlenie koloru niebieskiego, zaś temperatura końcowa, obliczona jako suma temperatury początkowej i przedziału temperatur, określa maksymalną temperaturę, dla której urządzenie prześle do diod LED składowe kolory RGB odpowiedzialne za wyświetlenie koloru czerwonego. Wartości spoza zakresu obcięte zostaną do wartości skrajnych, zaś wszystkie pośrednie powodują przesłanie składowych kolorów RGB obliczanych zgodnie z **rysunkiem 8**.

Nasuwa się pytanie, jak więc wprowadzamy niezbędne ustawienia? Urządzenie



Rysunek 8. Sposób wyznaczania składowych koloru diod LED na podstawie wartości mierzonej temperatury

po włączeniu zasilania przechodzi do trybu bezczynności, z wyłączonym wyświetlaczem LED (w celu oszczędzania energii). Jakikolwiek naciśnięcie przycisków funkcyjnych przełącza je w tryb regulacji temperatury początkowej (włączając

jednocześnie wyświetlacz LED) – w trybie tym krótkie naciśnięcie przycisków UP lub DOWN skutkuje zmianą wartości temperatury (w pętli). Długie naciśnięcie przycisku UP powoduje z kolei przejście do trybu regulacji przedziału temperatur,

w którym – jak poprzednio – krótkie naciśnięcie przycisków UP lub DOWN wywołuje zmianę bieżącej wartości (również w pętli). Długie naciśnięcie przycisku UP powoduje tym razem powrót urządzenia do trybu bezczynności i wyłączenie wyświetlacza z jednoczesnym zapisaniem ustawień w nieulotnej pamięci EEPROM mikrokontrolera, by mogły być wczytane jako startowe podczas włączania urządzenia. Co ważne, program wyposażono dodatkowo w zegar bezczynności, który po upływie 20 sekund braku reakcji ze strony użytkownika powoduje automatyczne przełączenie w tryb bezczynności i wyłączenie wyświetlacza, bez zapisywania dokonanych ustawień w pamięci EEPROM. Dodatkowo każdorazowo podczas dokonywania pomiaru temperatury (czyli co 1 sekundę) sprawdzana jest obecność czujnika temperatury i – w przypadku jego braku (lub problemów na magistrali 1-wire) – na wyświetlaczu widoczna jest informacja o błędzie w postaci litery „E” (dopóki dany problem nie ustąpi).

Robert Wołgajew, EP

REKLAMA

Wydawnictwo AVT nawiąże współpracę redakcyjną z osobami dobrze operującymi terminologią elektroniki i słowem pisanym. Propozycja szczególnie interesująca dla nauczycieli elektroniki, autorów artykułów, skryptów i książek.

**Aplikacje prosimy kierować na adres:
redakcja@elportal.pl**





Palniki i lutownice gazowe w ofercie TME

Palnik gazowy to jedno z mniej oczywistych narzędzi, które profesjonaliści z wielu branż powinni mieć w swoim wyposażeniu. Tymczasem gazowa lutownica stanowi wygodną, a czasem niezastąpioną alternatywę dla tradycyjnych, elektrycznych odpowiedników.

Przegląd podręcznych urządzeń zasilanych butanem

W niniejszym opisie prezentujemy cechy, zalety oraz przegląd palników i lutownic, w których źródłem energii jest spalanie butanu. W katalogu TME oferta wspomnianych urządzeń obejmuje kilkadziesiąt pozycji. Artykuły te różnią się pod względem rozmiarów i możliwości, ale również cen – chociaż należy już na wstępie podkreślić, że to nie koszt powinien decydować o doborze urządzenia. Wybór warto bowiem uzależnić przede wszystkim od potrzeb, indywidualnych preferencji oraz przewidywanej intensywności eksploatacji.

W naszym artykule poruszamy takie tematy, jak:

- typowe prace wykonywane z zastosowaniem palników gazowych,
- istotne cechy i właściwości palników oraz przydatne akcesoria,
- wybór palnika w zależności od przeznaczenia,
- najważniejsze zalety lutownic gazowych,
- zestawy akcesoriów dostarczane z lutownicami gazowymi,
- rozmiary i właściwości wyróżniające narzędzia zasilane butanem.

Palniki gazowe

Palniki gazowe w ofercie TME to niewielkie, podręczne przyrządy, których głównym przeznaczeniem są prace o stosunkowo małej skali. Przykładami takich zadań mogą być: **zacieśnianie**

tulei termokurczliwych, usuwanie niepożądanych powłok (lakierów/farb) z gwintów, podgrzewanie zabezpieczonych połączeń (hydraulika, mechanika samochodowa/rowerowa), osuszanie itp. Oczywiście można by tu wymienić również szereg innych dziedzin, w których stosowane są takie palniki, np. w laboratoriach chemicznych służą one do podgrzewania substancji, są też stosowane w kuchniach (do karmelizowania, nadpalania).

Podstawowe cechy

Główną cechą palnika gazowego jest zdolność do **długotrwałej emisji skoncentrowanego płomienia**. W przeciwieństwie do zwykłej zapalniczki, urządzenie nie przegrzewa się po kilkunastu sekundach, a jednocześnie pozwala szybko osiągnąć wysoką



Fotografia 1. Palnik WEL.WT13EU marki Weller

Fotografia 2. Palnik PORTA-GT220 marki Portasol

Fotografia 3. Palnik ARS-ES720TH marki Aries

temperaturę. W ofercie TME można znaleźć urządzenia tego typu o pojemności sięgającej **od 5,5 ml do 68 ml**, przystosowane do zasilania powszechnie dostępnym paliwem: **butanem (propanem-butanem)**. Wiele z nich należy do oferty znanych i wyspecjalizowanych marek, takich jak Weller, Portasol czy Aries.

Zabezpieczenia i rozwiązania konstrukcyjne

Niezależnie od producenta, palniki mają kilka cech wspólnych. Przede wszystkim są wyposażone w **zapalniki (piezoelektryczne)**, zawory regulacyjne oraz wbudowane zabezpieczenia uniemożliwiające przypadkowe załączenie. Niekiedy blokada zapłonu jest aktywowana automatycznie po nałożeniu **nasadki ochronnej**. Większość palników wyposażona jest w zintegrowaną lub nakładaną podstawkę, przez co urządzenie można postawić (chroniąc powierzchnię miejsca pracy przed kontaktem z rozgrzaną dyszą). Inna istotna cecha omawianych palników to **możliwość pracy w dowolnym położeniu**.



Fotografia 4. Palnik ROT-35130 marki Rothenberger Industrial

W ofercie TME są dostępne zarówno kompaktowe przyrządy o wysokości 72 mm, jak i większe (195 mm), które mogą służyć do **opalania czy podgrzewania obiektów o stosunkowo dużym rozmiarze** i dobrym przewodnictwie termicznym (np. podgrzewanie zabezpieczonych gwintów rur stalowych, drobne prace lutownicze instalacji miedzianych). Płomień emitowany przez palniki ma temperaturę **ok. 1300°C**, a moc urządzeń sięga **od 50 W do 850 W**.

Lutownice gazowe

Lutownice gazowe okazują się znakomitą alternatywą dla urządzeń elektrycznych, ponieważ do pracy nie wymagają źródła prądu, są przenośne (niewielkich rozmiarów i, siłą rzeczy, bezprzewodowe). Producenci wyposażenia lutowniczego dostarczają takie narzędzia w formatach pozwalających na wykonywanie **precyzyjnych prac lutowniczych, prostych napraw elektroniki, obróbki przewodów** etc. Dla wielu profesjonalistów będzie to znakomity dodatek do podręcznego ekwipunku, ułatwiający prace w terenie czy w nietypowym położeniu (przy naprawie maszyn i pojazdów).

Wyposażenie i rozmiary

Większość lutownic gazowych jest dostarczana z co najmniej kilkoma akcesoriami. Do najczęściej spotykanych należy skuwka,



Fotografia 5. Lutownica PORTAPRO-KIT marki Portasol

Poznaj ofertę kompaktowych lutownic i palników gazowych



Transfer Multisort Elektronik Sp. z o.o.
Łódź, Polska, dso@tme.pl

tme.eu

Dołącz do nas:





Fotografia 6. Lutownica DREMEL-VERSATIP marki Dremel

czyli nasadka zakrywająca dyszę. Urządzenia mogą mieć zestaw łatwo wymiennych grotów, w którym – oprócz elementów precyzyjnych (stożkowych) – najczęściej spotyka się również nożowe (o długiej, skośnej krawędzi mogącej służyć do cięcia tworzyw sztucznych) i groty o dużej powierzchni (zgrzewającej). Zróżnicowany wybór końcówek bywa istotny, ponieważ lutownica gazowa osiąga moc wyższą (nawet do 130 W) od urządzenia grzałkowego tych samych rozmiarów, a jej groty mogą rozgrzewać się do temperatury nawet 625°C (zależnie od modelu), w związku z czym zastosowania takich przyrządów są nieco szersze niż w przypadku wersji elektrycznych. Niektóre modele lutownic dostępne w ofercie TME mają w komplecie etui transportowe, służące jednocześnie jako organizator z miejscem na dodatkowe akcesoria (podstawkę/stojak, cyne, czyściwo etc.). Inne produkty odznaczają się jeszcze bardziej kompaktowymi wymiarami (np. 160 mm długości) i mogą być transportowane nie tylko w przegrodzie torby narzędziowej, ale nawet w kieszeni spodni czy koszuli.



Fotografia 7. Lutownica ROT-35120 marki Rothenberger

Zalety i wielofunkcyjność

Długość lutownic sięga od 160 mm do 210 mm. Osiągają one moce od 25 W do 130 W (przy czym w wielu modelach



Fotografia 8. Lutownica ERSA-OG13400141 marki Erska



Fotografia 9. Lutownica FUT.SKC-70 marki Engineer



Fotografia 10. Lutownica ARS-ES665P marki Aries

Fotografia 11. Lutownica JBC-SG1070 marki JBC Tools

przepływ gazu, moc i, co za tym idzie, temperatura grota mogą być regulowane). Dostępne w katalogu TME modele można napełniać standardowym butanem do zapalniczek (nie wymagają użycia specjalnego kartridża), dzięki czemu znalezienie i uzupełnienie paliwa nie stanowi problemu, jest też **bardzo tanie**. Oprócz atrakcyjnych parametrów, do istotnych cech urządzeń gazowych należy też **krótki czas nagrzewania grota** oraz (w przypadku niektórych artykułów, przy zastosowaniu specjalnej dyszy) możliwość pracy jako **lutownica na gorące powietrze**.

**Tekst opracowany przez
Transfer Multisort Elektronik Sp. z o.o.**

<https://www.tme.eu/pl/news/about-product/page/60334/oferta-palnikow-i-lutownic-gazowych/>



Fotografia 12. Lutownica WEL.PYROPEN-JUNIOR marki Weller



W ofercie AVT*

AVT6056

Najważniejsze parametry:

- generowanie rytmicznego dźwięku „bum-bum”,
- praca przez 15 min, 30 min lub 60 min,
- możliwość wcześniejszego wyłączenia układu,
- niski pobór prądu: około 100 nA w stanie spoczynku i nie więcej niż 1 mA podczas pracy,
- łatwa obsługa: wystarczy wcisnąć przycisk z wybranym czasem,
- zasilanie napięciem 3 V lub 4,5 V – dwie lub trzy baterie AA lub AAA.

* **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]** – 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

AVT5995	Elektroniczny gong (EP 8/2023)
-----	Sygnalizator ding-dong (EP 10/2018)
AVT793	Generator dźwięków alarmowych (EdW 11/2016, EP 12/2016)
AVT1897	Sterownik syreny piezo (EP 2/2016)
AVT1565	Elektroniczna syrena (EP 3/2010)
AVT1425	Miniaturowy sygnalizator alarmowy (EP 4/2006)
AVT2774	Syrena alarmowa dużej mocy (EdW 12/2005)
AVT740	Niezwykła „niebieska” dotykowa syrena policyjna.
	Uniwersalny generator VCO (EdW 10/2005)
AVT1304	Syrena z układem ZSD100 (EP 5/2001)

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.

Symulator dźwięku bijącego serca

Maluchy, w szczególności noworodki, uspokajają się, przebywając w otoczeniu jednostajnych dźwięków, które kojarzą im się dobrze. Takim dźwiękiem może być rytm serca ich mamy, którego słuchały przed porodem przez wiele tygodni. Prezentowany układ w prosty sposób wytwarza zbliżony dźwięk przez zadany czas.

Bum-bum, bum-bum, bum-bum... Każdy z nas zna ten odgłos. To dźwięk bijącego serca. Rytmiczny, jednostajny, monotonny, można wręcz powiedzieć: usypiający. Niemowlaki lubią takie dźwięki. Wsłuchują się w nie i zasypiają, śniąc swoje niemowlęce sny. Czemu by im tego nie ułatwić? Wszak wyspane dziecko to spokojniejsze dziecko – sprawdziłem tę zasadę w praktyce.

Dla współczesnej elektroniki wygenerowanie dźwięku, który choć trochę przypomina znane wszystkim „bum-bum”, to bułka z masłem. Oczywiście warto zadbać o bezpieczeństwo takiego urządzenia: powinno być ono zasilane niskim napięciem, z odseparowaniem od sieci elektrycznej. Nie powinno też pobierać zbyt dużo energii z baterii. Ani kosztować zbyt wiele. Mieszanka niemożliwa do zrealizowania? Wcale nie!

Budowa

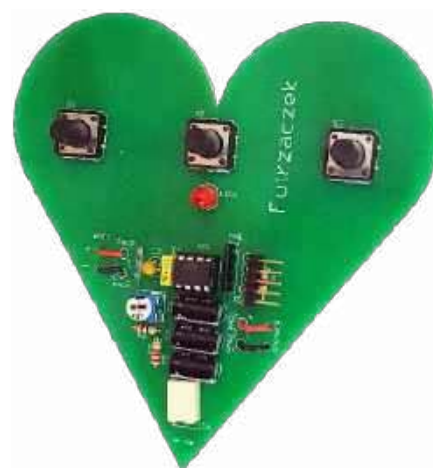
Schemat ideowy omawianego układu znajduje się na **rysunku 1**. Za generowanie impulsów odpowiada prosty, tani i dobrze znany mikrokontroler ATtiny13A. Przyciskami S1...S3 można zadać żądany czas wytwarzania dźwięku. Aby układ był bardziej odporny na zakłócenia

elektromagnetyczne, wejścia procesora, które sprawdzają stan styków tych przycisków, zostały dodatkowo podciągnięte rezystorami zawartymi w strukturze drabinki RN1. Dotyczy to również wejścia zerującego mikrokontroler.

Te wyprowadzenia MCU, które mogą posłużyć do jego zaprogramowania w układzie – bez wyciągania z podstawki – zostały wyprowadzone na złącze J1. Kondensatory C1 i C2 zmniejszają tętnienia napięcia zasilającego mikrokontroler, które mogą pojawiać się podczas jego pracy.

Dźwięk imitujący brzmienie bijącego serca jest generowany z unipolarnej fali prostokątnej, wytwarzanej przez mikrokontroler na jednym z wyjść cyfrowych. Szeregowe włączenie kondensatora C3 powoduje jego zróżniczkowanie i tym samym odcięcie drogi przepływu prądu dla składowej stałej, co zmniejsza zużycie energii z baterii. Szeregowe połączenie rezystora R1 i potencjometru P1 ustala głośność oraz uniemożliwia przepływ przez wyjście mikrokontrolera prądu o zbyt dużym natężeniu, w efekcie spowalniając przeładowywanie C3.

Równolegle do cewki głośnika, która powinna zostać podłączona do pól lutowniczych PAD1 i PAD2, zostały włączone



elementy tworzące filtr górnozaporowy. W dziedzinie wysokich częstotliwości blokadę tę stanowi kondensator C6, który zwiększa czas narastania napięcia na cewce głośnika. Tę samą funkcję pełni kondensator bipolarny o wypadkowej pojemności około 110 µF, który składa się z dwóch kondensatorów elektrolitycznych C4 i C5. Jednak wpływ tego kondensatora ulega zmniejszeniu, ponieważ jego szeregową impedancję zwiększa rezystor R2. Bez wspomnianych elementów filtracyjnych głośnik wydawałby z siebie bardzo głośne i nie naturalnie brzmiące trzaski.

Montaż i uruchomienie

Układ został zmontowany na jednostronnej płytce drukowanej w kształcie serca, której obrys mieści się w prostokącie o wymiarach 114 mm × 105,5 mm. Wzór jej

Wykaz elementów:

Rezystory: (THT o mocy 0,25 W)
R1, R2: 22 Ω
RN1: 4×10 kΩ SIL5
P1: 1 kΩ montażowy leżący

Kondensatory:

C1: 100 nF 63 V (raster 5 mm, MKT)
C2: 6,8 µF 10 V (raster 2,5 mm, tantalowy)

C3...C5: 220 µF 16 V (raster 2,5 mm)
C6: 1 µF 63 V (raster 5 mm, MKT)

Półprzewodniki:

U1: ATtiny13A (DIP8)

Pozostałe:

J1: goldpin 5 pin kątowy męski, 2,54 mm THT (opis

w tekście)

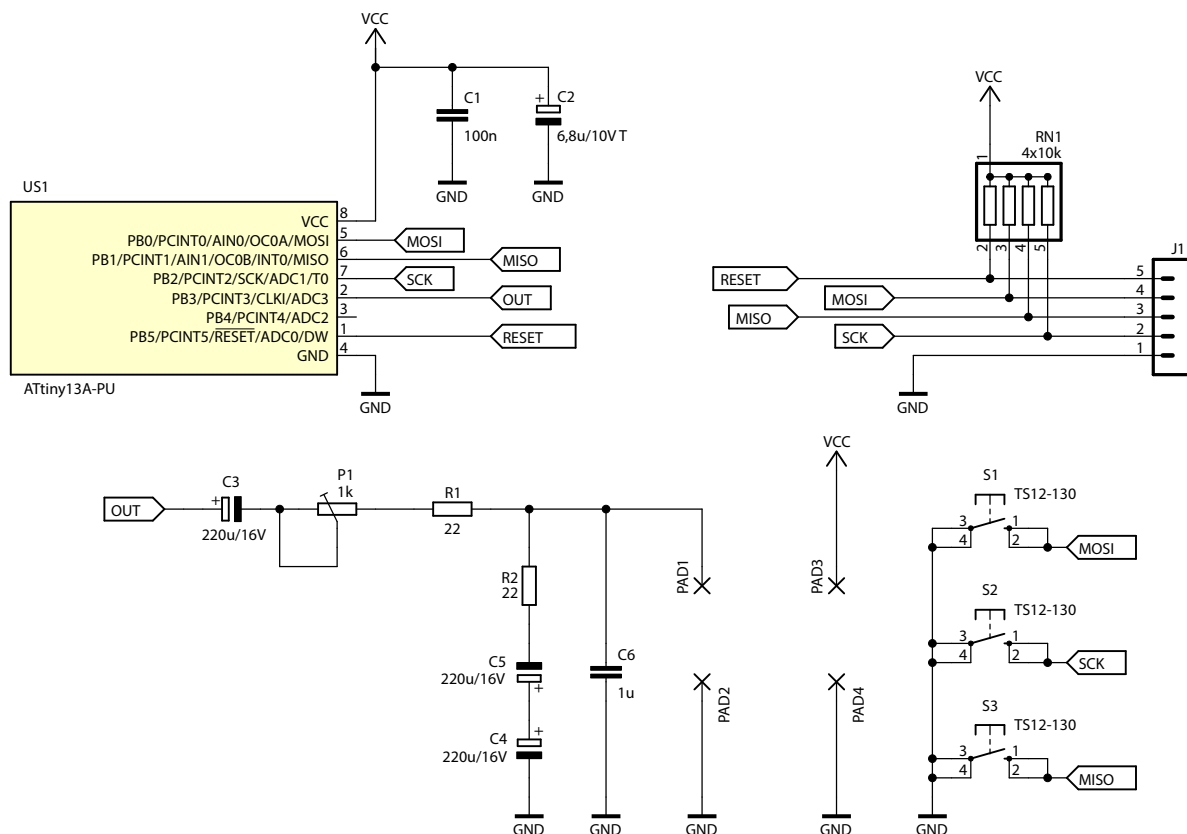
S1...S3: TS12-130

Jedna podstawka DIP8

Głośnik 8 Ω np. GŁ 8 Ω/1 W (opis w tekście)

Cienkie przewody połączeniowe

Koszyk baterii (opis w tekście)



Rysunek 1. Schemat ideowy układu symulatora dźwięku bijącego serca

ścieżek oraz schemat montażowy pokazano na **rysunku 2**.

Montaż najlepiej rozpocząć od elementów o najmniejszej wysokości obudowy,

czyli rezystorów R1 i R2. Pod mikrokontroler US1 proponuję zastosować podstawkę, aby ułatwić jego programowanie oraz wymianę w razie uszkodzenia. Kondensatory

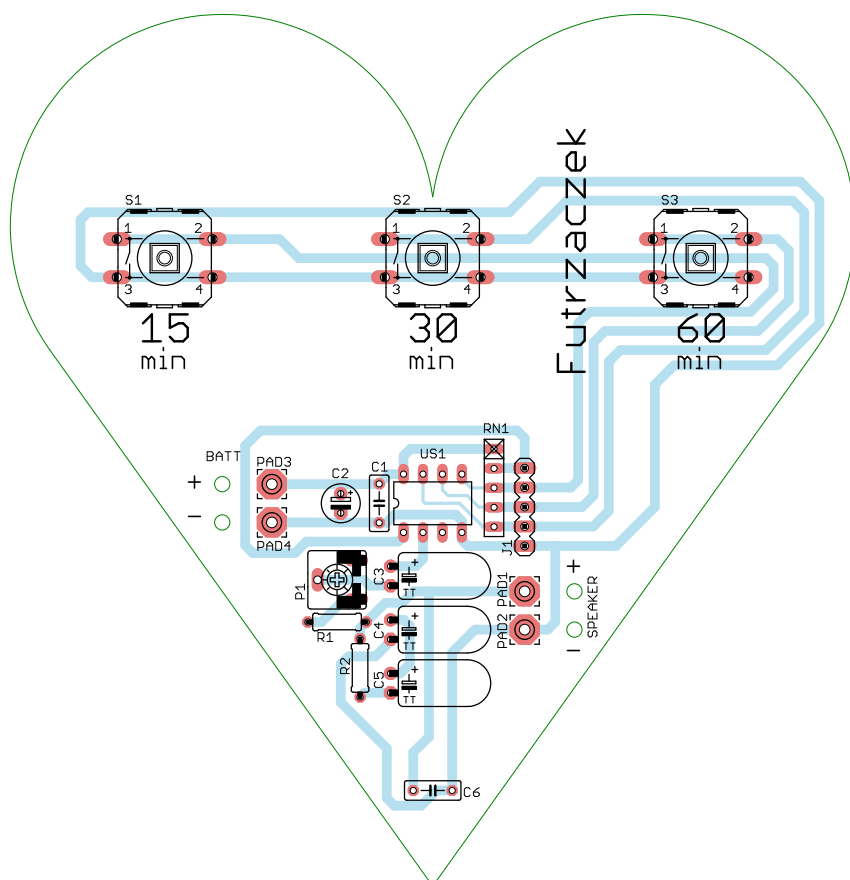
C3...C6 należy włutować na nóżkach na tyle długich, aby dało się je położyć na powierzchni laminatu. Złącza J1 można nie lutować, jeżeli mikrokontroler byłby programowany w zewnętrznym programatorze, wyposażonym we własną podstawkę. W pełni zmontowany układ prototypowy można zobaczyć na **fotografii tytułowej**.

Na etapie uruchamiania konieczne jest zaprogramowanie pamięci Flash mikrokontrolera dostarczonym wsadem. Wartości bitów zabezpieczających pozostają niezmienione względem ustawień fabrycznych i powinny wynosić:

Low Fuse = 0x6A

High Fuse = 0xFF

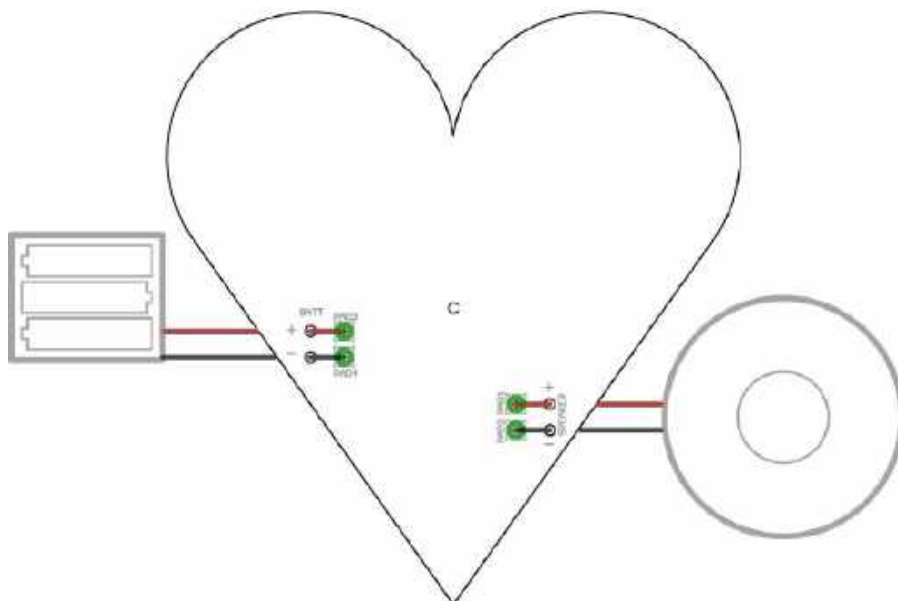
Szczegóły są widoczne na **rysunku 3**, który zawiera widok okna konfiguracji tychże bitów w programie BitBurner. Częstotliwość taktowania rdzenia mikrokontrolera będzie wynosiła zaledwie 1,2 MHz, co okazuje się wystarczające w tym zastosowaniu. Z kolei



Rysunek 2. Schemat montażowy i wzór ścieżek płytki



Rysunek 3. Konfiguracja bitów zabezpieczających



Rysunek 4. Schemat podłączenia układu

brak działania Brown-Out Detector przyczyni się do zwiększenia energooszczędności urządzenia.

Do płytki drukowanej należy podłączyć głośnik oraz koszyk baterii, które będą zasilaly układ. Najlepiej, aby głośnik miał impedancję $8\ \Omega$ i możliwie wysoką skuteczność, co przełoży się na wyższy poziom głośności, jaki można uzyskać. Dla przetwornika zostały przeznaczone pola lutownicze PAD1 i PAD2. Sam głośnik, po przylutowaniu cienkimi przewodami, można przymocować np. z tyłu płytki, pamiętając o tym, by metalowy kosz nie dokonał zwarcia na płycie.

Układ nadaje się do zasilania napięciem 3 V (dwie baterie AA lub AAA) lub 4,5 V (trzy ogniwa). Wyższe napięcie przekłada się na zwiększenie głośności maksymalnej układu – z tego powodu zalecałbym użycie koszyka na trzy baterie, podłączonego do pól lutowniczych PAD3 i PAD4. Taki koszyk nie wymaga wyłącznika i również można go przykleić z tyłu płytki, co dodatkowo ułatwiają otwory znajdujące się przy polach lutowniczych (pozwalające na przewleczenie przez nie przewodów poprowadzonych

z tylnej strony płytki, ponieważ zwiększa to wytrzymałość tego typu połączeń). Podłączenie jest widoczne na **rysunku 4**.

Układ został tak zaprojektowany, by pobierać jak najmniej energii z baterii, toteż przez większość czasu mikrokontroler pozostaje w stanie spoczynku. Pobór prądu w stanie czuwania nie przekracza 100 nA, natomiast podczas wydawania dźwięku wynosi średnio ok. 0,6 mA (przy napięciu 3 V) – lub ok. 0,8 mA (przy napięciu 4,5 V). Komplet baterii wystarczy zatem na wiele godzin „bum-bumania” przy dziecięcym uchu. Potencjometrem P1 można regulować głośność generowanych dźwięków.

Obsługa układu jest niezwykle prosta: po włożeniu baterii układ przechodzi w stan spoczynku. Kiedy wciśniemy jeden z przycisków S1...S3, układ uruchamia się na czas opisany pod przyciskiem, odpowiednio: 15 min, 30 min lub 60 min. Jeżeli w ciągu 15 s od włączenia nastąpi drugie wciśnięcie tego samego przycisku, układ wyłączy się. Jeżeli zaś upłynęło więcej czasu, to wciśnięcie przycisku spowoduje ponowne załadowanie licznika wartością odpowiadającą czasowi, na który wskazuje ów przycisk. Przykładowo:

wciśnięto S1 (15 min), po upływie 10 min okazało się, że dziecko dalej jest niespokojne, więc wciśnięto S3 (60 min). Układ będzie łącznie działał przez 70 min. Wyłączenie przed czasem również jest możliwe: wystarczy jeden z przycisków wcisnąć jeden raz (załadowanie pełnego czasu do odliczania), a zaraz potem wcisnąć go ponownie – układ potraktuje to zdarzenie tak samo, jakby został dopiero co wybudzony, nie zważając na fakt, że wcześniej już pracował.

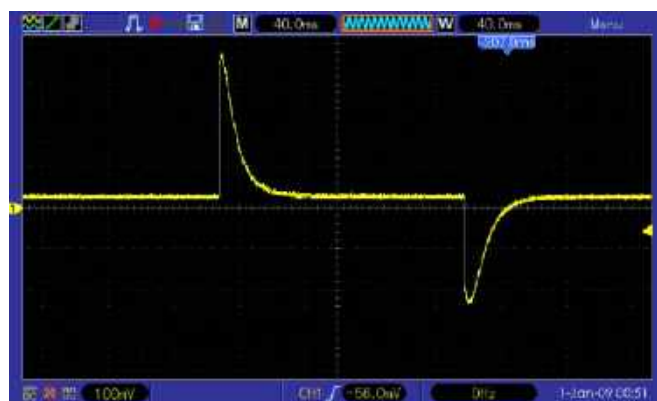
Wybudzenie mikrokontrolera ze stanu uśpienia jest możliwe po zwarcu styków S1...S3, ponieważ tak zostały skonfigurowane zewnętrzne przerwania. Jeśli opisana procedura przebiega prawidłowo, a do pamięci zostanie załadowany czas do odliczania, uruchomiony zostaje układ watchdog, który generuje przerwania co 125 ms. To one odpowiadają za cykliczne wybudzanie mikrokontrolera w trakcie generowania dźwięku. 8 takich przerwań to odliczenie pełnej sekundy. Po upływie całego zaprogramowanego czasu, watchdog jest wyłączany i mikrokontroler znowu przechodzi w stan pełnego uśpienia.

Opisany wcześniej filtr RC bierze udział w formowaniu impulsów pobudzających cewkę głośnika. Ciąg takich impulsów jest widoczny na **rysunku 5**. Na każdym osiem przerwań od watchdoga dwa ustawiają wyjście mikrokontrolera w stanie wysokim, zaś przez pozostałych sześć – jest ono w stanie niskim. Zbocze narastające owego sygnału daje impuls o polaryzacji dodatniej, zaś opadające – impuls o polaryzacji ujemnej. To tworzy charakterystyczne „bum-bum”, z wyraźnie dłuższą przerwą między tymi impulsami. Zbliżenie na wspomniane impulsy zawiera **rysunek 6**, na którym widać, że wierzchołki owych impulsów są zaokrąglone wskutek działania filtrów dolnoprzepustowych. Gdyby nie one, z głośnika wydobywałyby się donośne, drażniące trzaski.

Michał Kurzela, EP



Rysunek 5. Ciąg impulsów zasilających cewkę głośnika



Rysunek 6. Pojedynczy sygnał „bum-bum”

**Najważniejsze parametry:**

- zasilanie: zewnętrzny zasilacz USB-C PD 15 V/25 W (min.) lub ładowarka „laptopowa” 19...20 V,
- wyjścia: +3,3 V/1 A, +5 V/1 A, +12 V/1 A, -12 V/150 mA,
- wskaźniki zasilania: diody LED na wejściu oraz każdym z wyjść.

*** 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:

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

AVT6029	Zasilacz PoE 5/12 V 30 W (EP 3/2024)
----	Zasilacz do płytek stykowych ze złączem USB-C (EP 10/2023)
AVT5977	Stabilizator napięcia symetrycznego z regulacją współbieżną (EP 3/2023)
AVT5963	Zasilacz warsztatowy część 1 i 2 (EP 12/2022, 01/2023)
----	Modułowy zasilacz warsztatowy (EP 5/2022)
----	Regulowany zasilacz warsztatowy – RPS-02 (EP 4/2022)
AVT5915	Zasilacz 5 V/1 A z szerokim zakresem napięć wejściowych (EP 1/2022)
AVT5908	Beztransformatorowy impulsowy zasilacz sieciowy (EP 12/2021)
AVT5872	Regulowany zamiennik stabilizatora 78xx (EP 7/2021)
AVT1990	Regulowany zasilacz do płytek stykowych (EP 8/2018)

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.

W ofercie AVT*

AVT6057

Miniaturowy zasilacz napięć ustalonych

Opisany projekt jest prostym zasilaczem dostarczającym typowych napięć +3,3 V, +5 V oraz ± 12 V, często stosowanych przy uruchamianiu prototypów. Układ może być zasilany z łatwo dostępnego zasilacza laptopa o napięciu 19...20 V lub ładowarki USB-C PD z trybem 15 V, np. z oficjalnego zasilacza Raspberry Pi 5, o mocy >25 W. Maksymalna obciążalność wyjść 3,5/5/12 V wynosi 1 A, a w przypadku wyjścia -12 V wynosi 0,15 A.

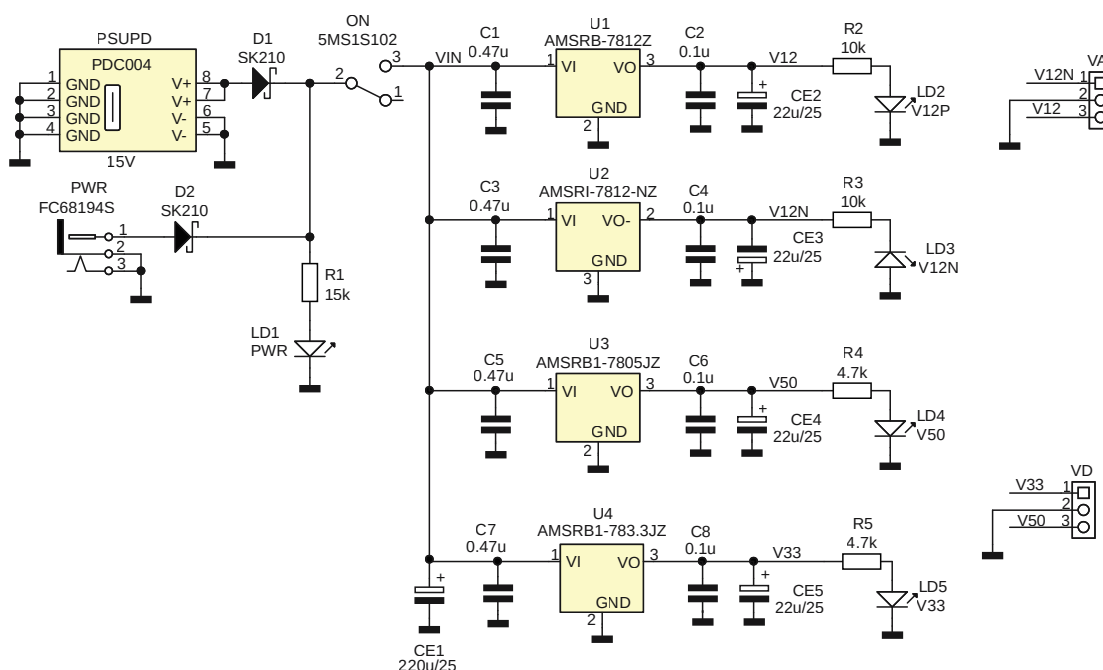


Schemat zasilacza pokazano na **rysunku 1**.

Układ oparty jest na scalonych przetwornicach U1...4 typu AMSRx o ustalonych napięciach wyjściowych. Dzięki zastosowaniu przetwornic możliwe jest uzyskanie

niewielkich wymiarów zasilacza oraz wysokiej sprawności (w porównaniu do klasycznych stabilizatorów 78xx/79xx). W celu uzyskania napięcia -12 V zastosowana została przetwornica obniżająco-odwracająca

AMSRL. Napięcie zasilania przetwornic może być doprowadzone do gniazda PWR 5,5/2,5 mm, a jego wartość powinna zawierać



Rysunek 1. Schemat ideowy zasilacza

Wykaz elementów:

Rezystory: (SMD 0805, 5%):

R1: 15 kΩ
R2, R3: 10 kΩ
R4, R5: 4,7 kΩ

Kondensatory:

C1, C3, C5, C7: 470 nF/25 V (SMD 0805, X7R)
C2, C4, C6, C8: 100 nF/25 V (SMD 0805, X7R)

CE1: 220 μF/25 V (6,3 mm elektrolityczny)
CE2...CE5: 22 μF/25 V (5 mm elektrolityczny)

Pozostałe:

PSUPD: moduł wyzwalacza USB 15 V (typ PDC004)
PWR: gniazdo zasilania 5,5/2,5 mm (FC68194S)
U1: AMSRL-7812Z (przetwornica obniżająca 12 V)
U2: AMSRL-7812-NZ (przetwornica obniżająca -12 V)

U3: AMSRB1-7805JZ (przetwornica obniżająca 5 V)
U4: AMSRB1-783.3JZ (przetwornica obniżająca 3,3 V)
ON: przełącznik suwakowy ON/OFF (typ 5MS1502)
VA, VD: złącze śrubowe 3,5 mm (typ DG381-3.5-3)

Półprzewodniki:

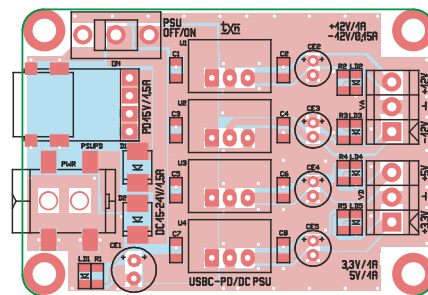
D1, D2: dioda Schottky'ego SK210 (SMB)
LD1...5: dioda LED (SMD 0805)

się w granicach 15...20 V – jako źródło zasilania układu może zatem posłużyć typowy zasilacz laptopa (wyposażony w odpowiedni adapter), który zalega zapewne w zakamarkach każdego warsztatu. Można w ten sposób dać drugie życie potencjalnemu elektrośmięciowi i stać się czynnym uczestnikiem ruchu „up-cyclingu”. Drugim sposobem zasilania jest zastosowanie typowej ładowarki USB-C PD o napięciu 15 V i mocy minimalnej 25 W. Aby ułatwić sobie zadanie, w zasilaczu użyto gotowego modułu wyzwalacza PD (o napięciu wyjściowym 15 V) typu PDC-004. Moduł jest łatwo dostępny na portalach aukcyjnych, w cenie kilku...kilkunastu PLN. Oczywiście można było zastosować któryś z układów wyzwalacza, np. IP2712 lub CHK22x, jednak ich cena i dostępność w ilościach jednostkowych jest dyskusyjna. Diody D1,2

zabezpieczają zasilacze przed zwarcie ich wyjść przy pomyłkowym równoczesnym podłączeniu obu źródeł. Wyłącznik ON umożliwia odłączenie zasilania przetwornic na czas manipulacji w zasilanym układzie. Diody LED LD1...5 sygnalizują obecność napięcia zasilania oraz napięć wyjściowych.

Układ zmontowano na dwustronnej płytce drukowanej, rozmieszczenie elementów pokazano na **rysunku 2**. Montaż jest typowy i nie wymaga dokładniejszego opisu. Moduł wyzwalacza USB-C – w celu zapewnienia dodatkowej stabilności mechanicznej – jest lutowany do padów SMD także w pobliżu gniazda USB-C. Zmontowany moduł zaprezentowano na **fotografii tytułowej**.

Poprawnie zmontowany układ nie wymaga uruchamiania. Należy sprawdzić jedynie



Rysunek 2. Rozmieszczenie elementów

poprawność lutowania i wyczyścić płytkę z resztek topnika. Po podłączeniu zasilania wystarczy skontrolować wartość napięć wyjściowych. W razie dłuższej pracy z pełnym obciążeniem warto na przetwornice nakleić niewielkie radiatory.

Adam Tatuś, EP

REKLAMA

Przejrzyj i zamów archiwalne wydania ELEKTRONIKI PRAKTYCZNEJ



PRZESYŁKA GRATIS



www.UlubionyKiosk.pl

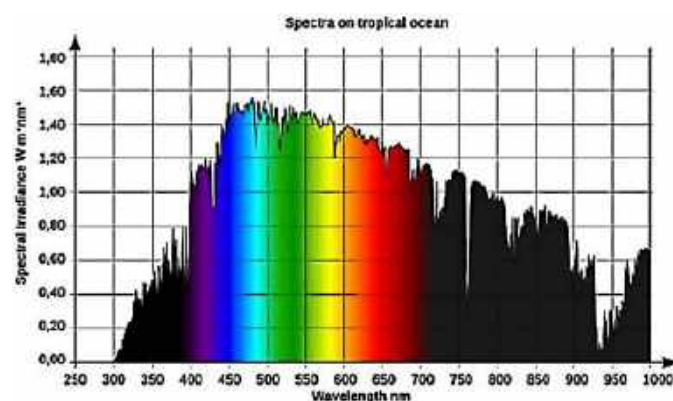
Parametry optyczne źródeł światła LED

Nowoczesna aparatura pomiarowa, stosowana przez producentów diod LED oraz bazujących na nich urządzeń, umożliwia twórcom dokładne scharakteryzowanie wszystkich istotnych wielkości fizycznych – określających zarówno elektryczne, jak i optyczne czy środowiskowe parametry źródeł światła. Dla elektronika, który z optyką ma do czynienia tylko okazjonalnie, spora liczba tabel i wykresów opisujących wszystkie aspekty emitowanego przez diody promieniowania może okazać się wręcz przytłaczająca – zwłaszcza jeżeli w urządzeniu mają być zastosowane diody białe, do których opisu stosuje się szczególnie dużo rozmaitych metod. Dokonajmy zatem szczegółowego przeglądu najważniejszych wielkości optycznych – rzecz jasna z punktu widzenia praktykującego konstruktora, którego zadaniem jest właściwy dobór źródeł światła do określonej aplikacji.

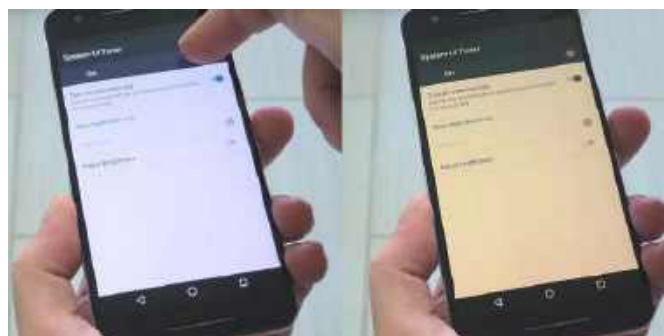
Jeśli nie wiadomo, o co chodzi...

... to chodzi o spektrum! I choć koszty źródeł światła LED mają istotny wpływ na praktyczny zakres ich zastosowań, to w najbardziej wymagających aplikacjach nawet finanse schodzą na dalszy plan. Dlatego naszą prezentację zaczniemy niejako „od końca”, skupiając się w pierwszej kolejności na złożonych aspektach związanych z parametrami i metodami porównywania widmowych charakterystyk źródeł światła.

Niedoścignionym wzorcem spektralnym dla współczesnej branży oświetleniowej jest naturalne światło słoneczne (rysunek 1). Próba jego odwzorowania ma znaczenie nie tylko ze względów estetycznych – liczne badania (prowadzone zarówno w dziedzinie fizjologii ludzi i zwierząt, jak i botaniki) wykazały ścisły związek długości fali z modulacją niektórych procesów biologicznych. Chyba każdy słyszał o niekorzystnym wpływie niebieskiej składowej światła (emitowanego przez urządzenia ekranowe) na nasz sen. Wynika to z faktu, że światło dzienne, bogate w fale o długościach rzędu 440...480 nm, na drodze milionów lat ewolucji stało się czynnikiem warunkującym aktywność komórek zwojowych siatkówki, przez co właśnie chłodne, niebieskawe światło najsilniej pobudza nasz organizm do aktywności. Z tego powodu – a także przez wzgląd na ryzyko rozwoju uszkodzeń siatkówki spowodowanych długotrwałą ekspozycją na niebieskie światło – współczesne urządzenia wyposażane są w funkcje „światła nocnego” lub „trybu czytania”,



Rysunek 1. Widmo światła słonecznego (<http://t.ly/hxrCs>)



Fotografia 1. Ilustracja działania trybu nocnego w smartfonie (<http://t.ly/HOURr>)



Fotografia 2. Pejzaż nadmorski zarejestrowany w „złotej godzinie” (<http://t.ly/piEL4>)

które ograniczają emisję światła właśnie z subpikseli „B” (fotografia 1). Mało tego: nawet wyższej klasy okulary produkuje się obecnie w wersji z filtrem o żółtawym zabarwieniu, który delikatnie ogranicza ekspozycję na niebieską część widma optycznego. Natomiast światło o ciepłym zabarwieniu (z większym udziałem składowych w górnym przedziale długości fali, czyli czerwień/żółć) wycisza i w naturalny sposób przygotowuje nasz mózg do nocnego spoczynku. I to także ma swoje źródło w ewolucyjnie wkodowanych wzorcach – wszak im bliżej zachodu słońca, tym wyraźniej rozkład widma światła słonecznego (przechodzącego przez kolejne warstwy atmosfery pod coraz większym kątem) przesuwają się w stronę barw cieplejszych (stąd też chwile tuż przed zniknięciem słońca za horyzontem, a – gwoździ ścisłości – także tuż po jego wschodzie, noszą w fotografii pejzażowej miano „złotej godziny”, czyli „golden hour” – fotografia 2).

Poznanie i zrozumienie wpływu światła na samopoczucie człowieka jest podstawą nie tylko



Fotografia 3. Żarówka LED marki Philips o barwie światła przestrajalnej za pośrednictwem Wi-Fi (<http://t.ly/jwp7e>)



Fotografia 4. Zastosowanie specjalnych lamp LED do stymulacji wzrostu roślin doniczkowych (<http://t.ly/VqnGc>)

rozmaitych technik fototerapeutycznych, ale także... tzw. „oświetlenia cyrkadialnego”, czyli mającego na celu wsparcie osób, przebywających w pomieszczeniach ze sztucznym oświetleniem, w efektywnej pracy i skutecznym wypoczynku. Coraz większą popularność zdobywają bowiem specjalne lampy odwzorowujące widmo światła słonecznego, ba! pozwalające na dynamiczną zmianę barwy światła w zaskakująco szerokim zakresie (**fotografia 3**).

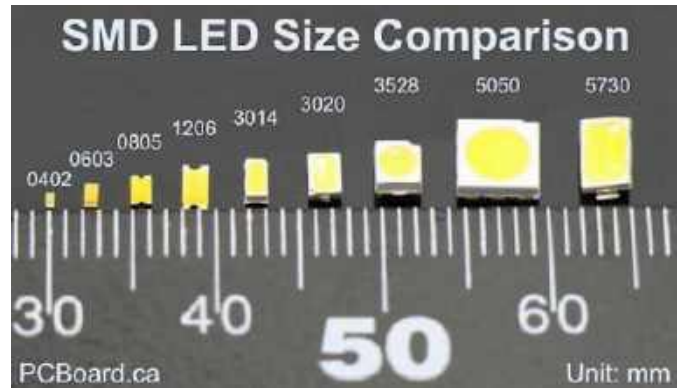
Przykłady z dziedziny fotobiologii można oczywiście mnożyć – długość światła moduluje wszak także procesy wzrostu roślin, a nawet sam chlorofil ma ściśle określoną charakterystykę spektralną (i to różną w zależności od odmiany tego związku). Stąd też wynika popularność, jaką w ostatnich latach zdobyły niedrogi lampy LED przeznaczone właśnie do sztucznego „napędzania” rozwoju roślin doniczkowych (i nie tylko) – patrz **fotografia 4**.

Choć postęp, który dokonał się na przestrzeni ostatnich trzech dekad w dziedzinie białych źródeł światła LED, jest naprawdę fascynujący, to okazuje się, że założenia leżące u podstaw tej technologii działają w istocie... na jej niekorzyść! I zdecydowanie nie chodzi tutaj o przysłowiowego „kolosa na glinianych nogach” – nic bowiem nie wskazuje na to, by dla białych LED-ów rośla na horyzoncie jakiegokolwiek znacząca konkurencja w dziedzinie iluminacji wszystkich, co tylko człowiek potrzebuje oświetlić. Problem leży głęboko w strukturze diod, a metody jego zwalczania są przedmiotem wyścigu technologicznego i to zarówno w samej branży optoelektronicznej, jak i we wszystkich branżach od niej zależnych – m.in. w dziedzinie sprzętu medycznego, fotografii, kinematografii czy wreszcie... konsumenckiego i przemysłowego oświetlenia pomieszczeń. Przyjrzyjmy się zatem bliżej tej tematyce, jako że jej zrozumienie pozwala znacznie pewniej poruszać się w zawiłościach współczesnej optoelektroniki i branży oświetleniowej.

Metody wytwarzania światła białego

Dzisiejsza branża oświetleniowa bazuje w znakomitej większości na białych diodach LED, które w istocie stanowią... modyfikację diod niebieskich. Gwoli ścisłości należy dodać, że istnieją obecnie trzy techniki wytwarzania światła białego za pomocą emiterów półprzewodnikowych.

- 1. Diody z powłoką fosforową** – zdecydowanie najpopularniejsza konstrukcja białych diod LED bazuje na strukturze niebieskiej lub (rzadziej) ultrafioletowej, której podstawowym budulcem jest azotek indowo-galowy (InGaN). Struktura ta jest w procesie produkcji pokrywana specjalnym luminoforem (**fotografia 5**), który – pobudzony światłem o stosunkowo małej



Fotografia 5. Białe diody LED SMD z widoczną warstwą żółtego luminoforu (<http://t.ly/vGz4o>)

długości fali – wytwarza szerokopasmowe promieniowanie przesunięte w stronę czerwieni. Zestawienie podstawowego pików światła niebieskiego z „rozlanym” pikiem pochodzącym z luminescencji powoduje wytworzenie światła białego.

- 2. Konstrukcje z mieszaniem barw składowych** znajdują zastosowanie wszędzie tam, gdzie emisja światła białego stanowi tylko jeden z trybów pracy. Jeżeli bowiem wymagane jest także uzyskiwanie „dowolnych” innych kolorów z szerokiej palety barw, to nie ma innej drogi jak tylko mieszanie światła w trzech pasmach podstawowych (czerwony, zielony, niebieski), w odpowiednich proporcjach. Jak powszechnie wiadomo, dokładnie tak działają wszystkie stosowane w praktyce diody RGB (**fotografia 6**), moduły, taśmy czy żarówki RGB, nie wspominając o wszelkiego rodzaju ekranach LED (np. telebimach), OLED, AMOLED, etc.



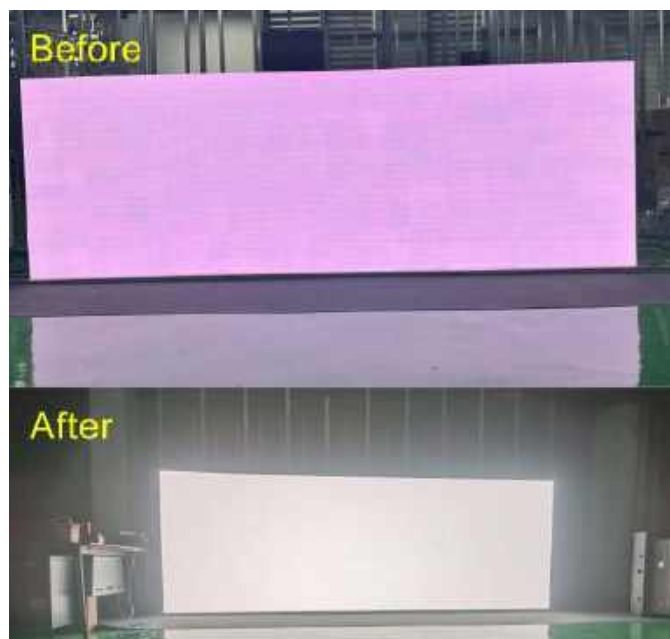
Fotografia 6. Przykładowa dioda LED RGB (<http://t.ly/KK2bE>)

- 3. Konstrukcje hybrydowe** – rzadziej spotykane jest zastosowanie rozwiązania hybrydowego, w którym diody o budowie opisanej w punkcie 1 są niejako „uzupełniane” przez starannie wyselekcjonowane struktury monochromatyczne. Ich zadaniem jest odpowiednie dostrojenie spektrum poprzez podcięcie go w pasmach niejako „ominiętych” przez proste diody z luminoforem. Szczególną odmianą tego rozwiązania stanowią diody RGBW (**fotografia 7**), powstające przez uzupełnienie klasycznej diody trójkolorowej o dodatkową, czwartą strukturę szerokopasmową, wykonaną rzecz jasna jako wspomniana wcześniej dioda z powłoką fosforową.



Fotografia 7. Dioda LED RGBW dużej mocy z serii XLamp® XM-L® Color Gen 2 (<http://t.ly/DdjHO>)

Tego rodzaju diody występują m.in. w wysokiej jakości ekranach LED, w których dokładne odwzorowanie bieli jest szczególnie istotne – należy bowiem pamiętać, że zmieszanie trzech, stosunkowo wąskich przedziałów widma (R, G, B) powoduje powstanie swego rodzaju „luk”, co znacząco pogarsza odwzorowanie delikatnych przejść tonalnych. Co gorsza: drobne rozrzuty produkcyjne (zarówno pomiędzy samymi diodami, jak i poszczególnymi kanałami sterowników prądowych) wpływają na powstawanie różnic tonalnych



Fotografia 8. Ekran LED przed i po kalibracji barwnej (<http://t.ly/tYXou>)



Fotografia 9. Wymiana modułu telebimu LED (<http://t.ly/KkUcu>)

powodujących wręcz fatalne wrażenia wizualne – co doskonale widać na **fotografii 8**, na której już na pierwszy rzut oka widać granice pomiędzy poszczególnymi modułami współtworzącymi panoramiczny telebim. Kalibracja jest oczywiście wymagana nie tylko na etapie instalacji, ale także po wymianie modułu (**fotografia 9**) – chyba każdemu zdarzyło się oglądać ekran reklamowy czy LED-owe tło na scenie koncertowej, w którym jeden lub kilka „prostokątów” wyraźnie odróżniało się od pozostałych pod względem kolorystycznym. Warto przy okazji dodać, że technologia RGBW pozwala nie tylko poprawić odwzorowanie barw, ale także... obniżyć pobór energii i uprościć sterowanie. Komercyjny sukces pierwszych diod RGBW spowodował, że na rynku pojawiły się nawet diody RGBCW, czyli... 5-kanalowe, które oprócz trzech kanałów podstawowych (RGB) zawierają także dwie białe struktury LED różniące się temperaturą barwową (**fotografia 10**).

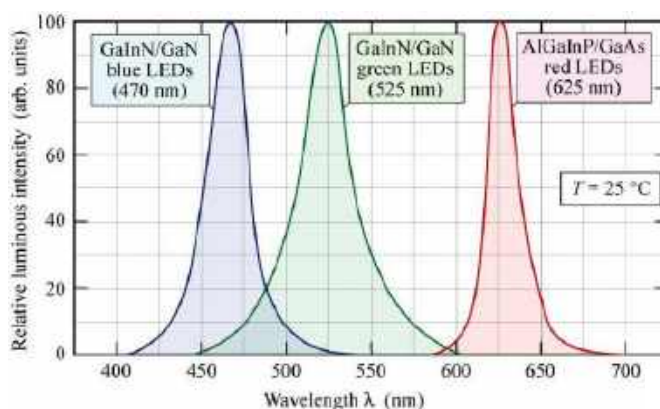


Fotografia 10. Dioda RGBCW (<http://t.ly/i8D2D>)

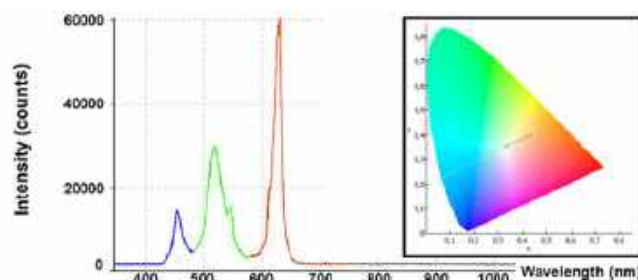
Jak widać, mamy pod ręką całkiem pokazny arsenał technik umożliwiających wytwarzanie światła białego. Czemu zatem na rynku wciąż pojawiają się nowe konstrukcje zarówno samych emiterów, jak i wszelkiego rodzaju diodowych lamp, modułów itp.? Odpowiedź na to pytanie wymaga pochylenia się nad szczegółami spektralnymi sztucznego światła białego.

Co w widmie błyszczy, czyli o tym, że białe białemu nierówne

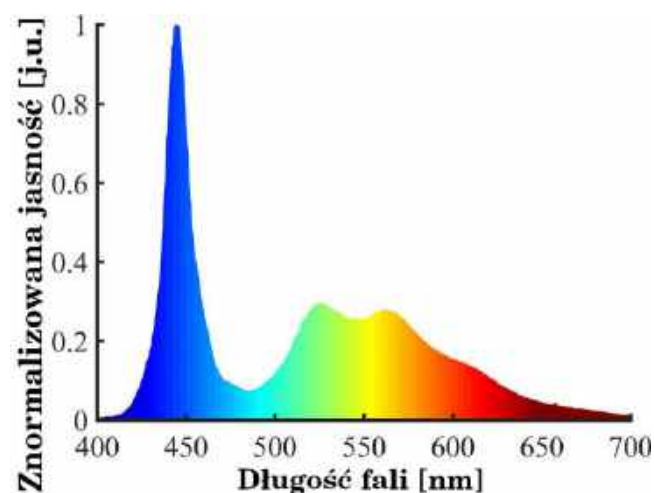
W praktyce okazuje się, że nawet pozornie „czyste”, neutralne światło białe wcale nie jest „prawdziwym” światłem białym. Wspomnieliśmy już, że niedoścignionym wzorem dla twórców współczesnych systemów oświetleniowych i komponentów optoelektronicznych jest naturalne, szerokopasmowe światło słoneczne. Jego widmo pokazaliśmy już na rysunku 1. Próba odwzorowania go za pomocą diod RGB musi rzecz jasna spalić na panewce – wystarczy spojrzeć na wykres pokazany na **rysunku 2**, a obrazujący typowe widma emisyjne struktur RGB. Oczywiście – zmieszanie ich w odpowiednich proporcjach sprawi wrażenie



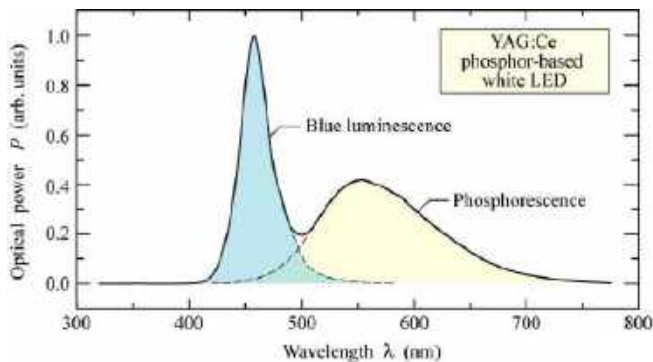
Rysunek 2. Typowe widma emisyjne diod LED w trzech kolorach podstawowych (<http://t.ly/LUP6G>)



Rysunek 3. Widmo rzeczywistej diody LED RGB świecącej kolorem białym o temperaturze barwowej 5507 K (<http://t.ly/JDPHP>)



Rysunek 4. Widmo typowej, białej diody LED (<http://t.ly/kXtff>)



Rysunek 5. Wyjaśnienie kształtu widma typowych diod białych (<http://t.ly/LUP6G>)

światła ciepłego, neutralnego lub zimnego (przykład na **rysunku 3**), ale z naukowego (a jak się za chwilę przekonamy – także użytkowego) punktu widzenia do ideału będzie jeszcze bardzo, bardzo daleko.

Jak zatem prezentują się pod tym względem klasyczne diody białe? Zdecydowanie lepiej, ale wciąż jeszcze nieidealnie. Widmo typowej diody z powłoką fosforową pokazano na **rysunku 4**, a wyjaśnienie jego genezy – na **rysunku 5**. Uwagę zwracają tutaj dwa aspekty. Pierwszym, który najbardziej rzuca się w oczy, jest znacznie lepsze „wypełnienie” widma w zakresie fal o średniej i dużej długości fali (rzecz jasna odnosimy się tutaj do granic pasma widzialnego). Drugą, doskonale widoczną cechą charakterystyczną zaprezentowanego na **rysunku 4** spektrum jest silny pik w rejonie barwy niebieskiej – nietrudno się domyślić, że jest za niego odpowiedzialna struktura diody pobudzającej żółto-pomarańczowy luminofor. Niestety, pomiędzy pikiem światła niebieskiego a resztą widma widoczna jest „dolina” – określana fachowo mianem *cyan gap*. I właśnie te dysproporcje w rozkładzie spektrum stanowią największy problem współczesnych, białych LEDów. Nietrudno się domyślić, że część światła niebieskiego przechodzi siłą rzeczy pomiędzy cząstkami luminoforu, a więc niełatwo jest osłabić to zjawisko: dołożenie jeszcze grubszej warstwy luminoforu mogłoby pogorszyć efektywność świecenia, z kolei obniżenie mocy zasilania diody doprowadzi do ogólnego spadku jasności we wszystkich zakresach widma. Co można z tym zrobić?

Metod jest, jak zwykle, kilka. Po pierwsze: można zastosować metodę hybrydową – dołożyć do białych struktur dodatkowe diody monochromatyczne, które uzupełnią widmo tylko „tam, gdzie trzeba”. Po drugie: istnieje możliwość przesunięcia piku emisji struktury podstawowej np. w stronę fioletu, co robi się już w niektórych przypadkach. Po trzecie: można modyfikować materiał samego luminoforu. Rzecz jasna, każda z metod ma swoje zalety, ale i pewne ograniczenia. Jedno jest pewne – otrzymanie dokładnej „kopii” światła słonecznego wydaje się na ten moment praktycznie niemożliwe, a nawet jeżeli – to bardzo trudne. Jak zatem zmierzyć, na ile udało nam się zbliżyć do ideału? Z pomocą przychodzi szereg technik spektrometrycznych: pozwalają one uzyskać zestaw standaryzowanych parametrów widmowych umożliwiających porównywanie poszczególnych źródeł światła.

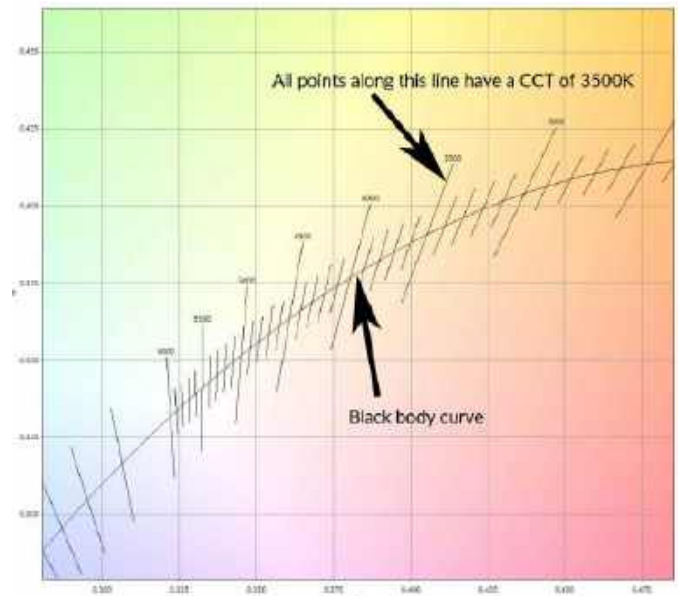
Parametry spektralne źródeł światła LED

Temperatura barwowa i CCT

Zdecydowanie najbardziej rozpowszechnioną metodą określania barwowego charakteru światła białego jest powszechnie znana temperatura barwowa (fotografia 11), u podstaw której leży porównanie danego odcienia światła (w przybliżeniu) białego z promieniowaniem termicznym ciała doskonale czarnego, rozgrzanego do danej temperatury. Stąd też jednostka temperatury barwowej, czyli Kelwin (K). Co ciekawe, z technicznego punktu widzenia pojęcie temperatury barwowej powinno być poprawnie stosowane



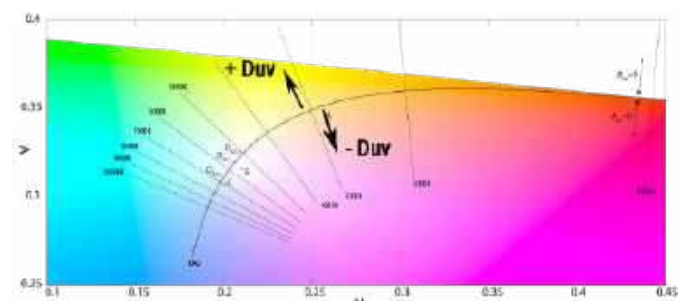
Fotografia 11. Porównanie źródeł światła o temperaturze barwowej od 1000 do 10000 K (<http://t.ly/5cwG7>)



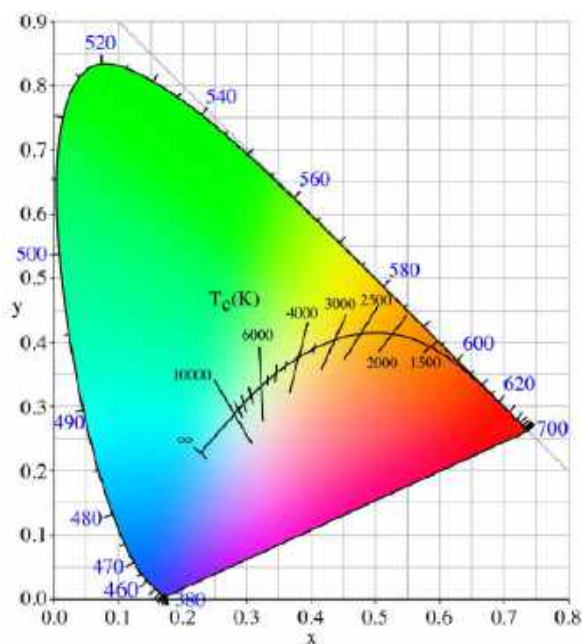
Rysunek 6. Krzywa promieniowania ciała doskonale czarnego z naniesionymi oznaczeniami odcinków odpowiadających poszczególnym wartościom CCT (<http://t.ly/a6gK1>)

wyłącznie w odniesieniu do punktów leżących bezpośrednio na krzywej promieniowania ciała doskonale czarnego (tzw. *black body curve*), a przecież wiemy już, że żadne sztuczne źródło światła nie jest idealne, zaś na finalną barwę wpływa szereg różnych czynników konstrukcyjnych, materiałowych i aplikacyjnych. Co zatem zrobić w tej sytuacji?

Możliwości jest, jak zwykle, wiele. W dokumentacjach elementów optoelektronicznych i gotowych źródeł światła (np. żarówek LED) często można natknąć się na skrót CCT, oznaczający – w dosłownym tłumaczeniu – skorelowaną temperaturę barwową. Pojęcie „skorelowany” wiąże się z faktem matematycznej próby objęcia – za pomocą jednej, zunifikowanej miary – nie tylko barw „leżących na krzywej”, lecz także nieco od niej odbiegających. Na **rysunku 6** można zobaczyć zestaw linii, które – przecinając krzywą odpowiadającą promieniowaniu ciała doskonale czarnego – wkraczają w obszary bardziej odległe od tego, co powszechnym rozumieniu nazwalibyśmy słowem „biel”. Taka metoda określania barw rodzi jednak następny problem: wartość CCT (także wyrażana w Kelwinach) jest



Rysunek 7. Sposób oznaczania wartości Duv (Delta u,v) w przestrzeni barwowej CIE 1960 UCS (<http://t.ly/FACVD>)



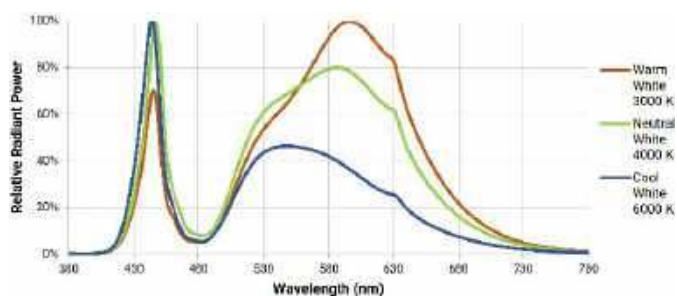
Rysunek 8. Przestrzeń barwna CIE 1931 (CIE XY) z oznaczeniem krzywej ciała doskonale czarnego i odcinków CCT (<http://t.ly/FACVD>)

wysokość niejednoznaczna, a zatem przydałby się sposób na scharakteryzowanie odległości danego punktu w przestrzeni barwnej od naszej „idealnej” krzywej odniesienia. Tak powstała miara Duv (**rysunek 7**), której celem jest właśnie umiejscowienie określonego koloru – pod względem jego odcienia i nasycenia – względem wspomnianej już krzywej. Dla lepszego zobrazowania omawianego zagadnienia spojrzmy jeszcze na popularny wykres przestrzeni barw wg standardu CIE 1931 (**rysunek 8**) – tutaj także znajdujemy odniesienie do krzywej promieniowania ciała doskonale czarnego wraz z zaznaczeniem poszczególnych odcinków odpowiadających skorelowanym temperaturom barwowym (CCT).

Współczynnik renderowania barw

Nie ulega wątpliwości, że temperatura barwowa i CCT to bardzo wygodne i użyteczne, skalarne miary odcienia światła białego, zwłaszcza w przypadku oświetlenia pomieszczeń. Niestety, z technicznego punktu widzenia wartości te nie mówią praktycznie nic o kwestii odwzorowania koloru powierzchni, na które pada światło o danej temperaturze barwowej.

I tutaj pojawia się spory problem. Jeżeli bowiem weźmiemy pod uwagę, że dana wartość CCT odnosi się do pewnego wycinka palety barw, a każde źródło sztuczne emituje widmo optyczne składające się z zestawu fal o różnych długościach i natężeniach, to dość szybko dojdziemy do wniosku, że pewne kolory powierzchni mogą być gorzej widoczne przy danym zestawieniu składowych, a lepiej – przy oświetleniu wiązką o innym składzie widmowym. W celu uzyskania pełniejszego obrazu sytuacji spojrzmy na charakterystyki trzech modeli diod dużej mocy z serii XLamp® XM-L®



Rysunek 9. Porównanie charakterystyk widmowych trzech modeli diod RGBW dużej mocy marki Cree (dot. tylko struktury białej) – <http://t.ly/OrFch>

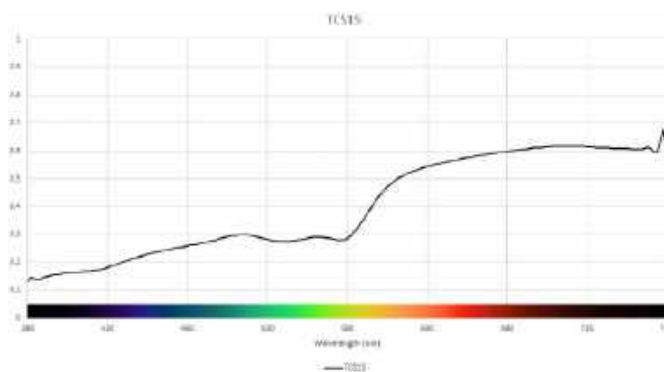


Rysunek 10. Paleta barw testowych stosowana do określenia CRI (<http://t.ly/fWY72>)

Color Gen 2 marki Cree, pokazane na **rysunku 9** – jak widać, poszczególne widma różnią się nie tylko stosunkiem wysokości pików światła niebieskiego do pików emisji luminoforu, ale także kształtem tego drugiego. I właśnie ze względu na wspomniane ograniczenia metod oceny CCT powstała ustandaryzowana miara określana mianem współczynnika renderowania barw (Color Rendering Index – CRI). Korzysta ona z zestawu specjalnie dobranych i dokładnie określonych, piętnastu kolorów testowych (**rysunek 10**), przy czym w pomiarach spektrometrycznych do każdego z nich jest przypisywana wartość od 0 do 100 – im wyższa, tym lepiej badane źródło światła radzi sobie z odwzorowaniem owego koloru. W szczególnym przypadku CRI może wynosić zero (światło monochromatyczne, poprawnie oddające tylko jeden określony kolor) lub 100 (światło słoneczne, a także żarowe czy halogenowe), zaś dla typowych, białych diod LED wartość ta wynosi przeważnie od 80 do 90.

I tutaj pojawia się kolejny problem. Wartość CRI, którą znajdujemy w większości not katalogowych czy opisów produktów, bazuje na uśrednieniu zestawu tylko... 8 „słupków” (czyli indeksów R1...R8) spośród wszystkich 15 próbek określonych normą. A co z resztą?

Niestety, tu właśnie leży pies pogrzebany. Klasyczna miara CRI (określana też mianem Ra od ang. „average”) bazuje wyłącznie na kolorach pastelowych. Tymczasem często równie ważne jest odwzorowanie barw nasyconych (R9...R12) oraz tzw. „kolorów ziemi” (R13...R15). Dlatego też w niektórych przypadkach producenci specyfikują także pozostałych siedem indeksów (R9...15), co stanowić ma jednoznaczny dowód, że oferowane przez nich źródła światła



Rysunek 11. Spektrum odcienia skóry osób pochodzenia azjatyckiego (<http://t.ly/fWY72>)



Fotografia 12. Klasyczny przykład obrazujący różnice w percepcji kolorów przy oświetleniu tego samego obiektu za pomocą źródeł o różnych wartościach CRI (<http://t.ly/Fhs74>)

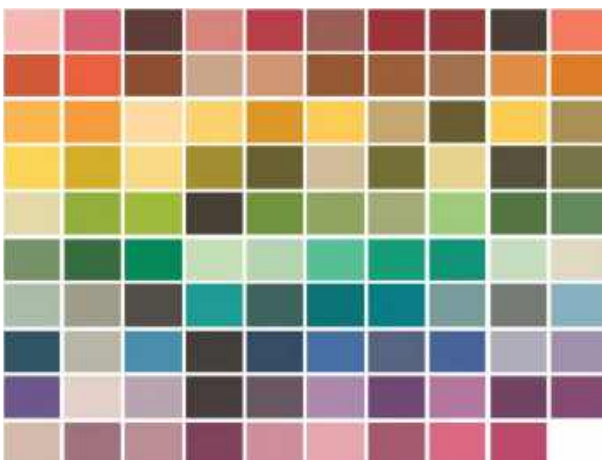
odzworowują naturalne spektrum światła słonecznego najlepiej, jak to tylko możliwe – i to pod każdym względem (czytaj – w pełnym zakresie widma widzialnego). W profesjonalnych aplikacjach (oświetlenie medyczne, fotograficzne czy kinematograficzne) szczególnie ważne jest prawidłowe oddanie nasyconej barwy czerwonej (R9), a także odcieni skóry – R13 i R15. Przykładowe spektrum barwy nr 15 (odcieni skóry osób pochodzenia azjatyckiego) pokazano na **rysunku 11**, zaś klasyczny przykład pozwalający na zaobserwowanie różnic pomiędzy widokiem tego samego obiektu oświetlonego źródłami o różnych wartościach CRI zaprezentowano na **fotografii 12**.

Standard TM-30-15/18

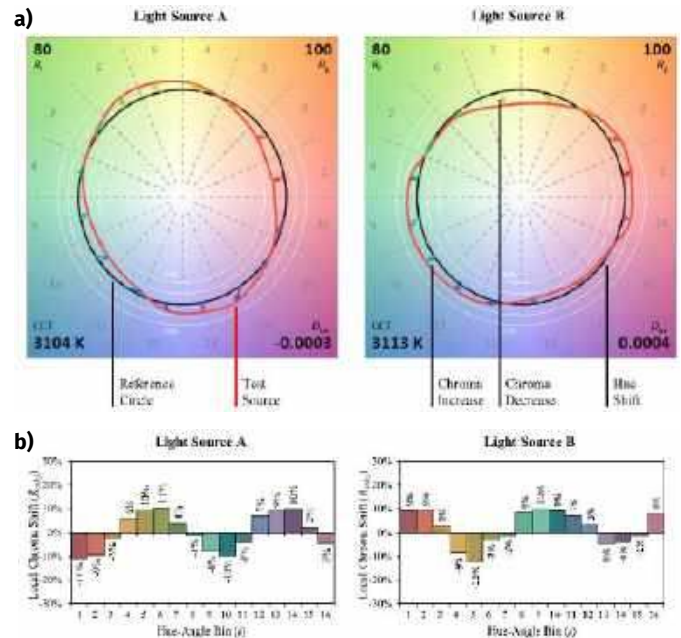
Uważny Czytelnik zapewne zada w tym momencie pytanie: a dlaczego akurat 15 kolorów testowych i dlaczego akurat te, a nie inne? Pytanie jest o tyle zasadne, że w (teoretycznie) nieskończonej liczbie możliwości z pewnością dałoby się wskazać także inne zestawy barw, które lepiej sprawdziłyby się w pewnych specyficznych sytuacjach. Warto wspomnieć chociażby o pierwszym z brzegu przykładzie aplikacji: lampy oświetlające ciała modelek w czasie profesjonalnych sesji fotograficznych czy też aktorów na planach hollywoodzkich superprodukcji muszą oddawać możliwie jak najlepiej subtelności koloru ludzkiej skóry. A tych, jak wiemy, jest bardzo wiele – nawet w ramach każdej rasy istnieje spora zmienność osobnicza w tym aspekcie.

I właśnie na bazie tego rodzaju wątpliwości (a także coraz powszechniejszej krytyki przydatności CRI w zastosowaniach profesjonalnych), w 2015 roku powstała norma stosowana do oceny kolorystyki światła białego, określana niewiele mówiącym skrótem TM-30-15. Trzy lata później została ona zaktualizowana do postaci TM-30-18. W obydwu przypadkach do określenia wierności odwzorowania barw stosuje się już nie 15, ale aż 99 kolorów testowych (**rysunek 12**), które... wydają się naturalne, przyjemne w percepcji i generalnie rzecz ujmując „bliskie”, prawda? Nic dziwnego – zostały one wybrane na podstawie analizy spektralnych pomiarów refleksyjności ponad 100 tysięcy (!) obiektów opisanych w ogólnościowych zasobach literatury naukowej – ze szczególnym uwzględnieniem kolorów skóry czy też kwiatów. Najpierw zdecydowano się na wybór 4880 próbek, a dopiero z nich utworzono reprezentatywny zbiór wspomnianych już 99 „barw świata”.

Mało tego – aby jeszcze mocniej uprościć korzystanie z wyników pomiarów, wprowadzono specjalny rodzaj wykresu polarnego, określony mianem Color Vector Graphics (CVG) – wszystkich 99 próbek jest bowiem łączonych w 16-słukowy histogram, a następnie reprezentowanych poprzez zamknięty, krzywoliniowy obrys na tle koła referencyjnego (**rysunek 13a**). Warto przeanalizować odpowiadające obrazom CVG wykresy słupkowe (**rysunek 13b**), które reprezentują przesunięcia chromatyczne w poszczególnych wycinkach spośród 16 równo rozłożonych obszarów przestrzeni barw.



Rysunek 12. Kolory testowe wg normy TM-30-15 (<http://t.ly/deHWH>)



Rysunek 13. Graficzne reprezentacje odwzorowania barw za pomocą dwóch różnych źródeł światła, przedstawione w formacie CVG (a) oraz odpowiadające im wykresy słupkowe (b) – <http://t.ly/4m4jc>

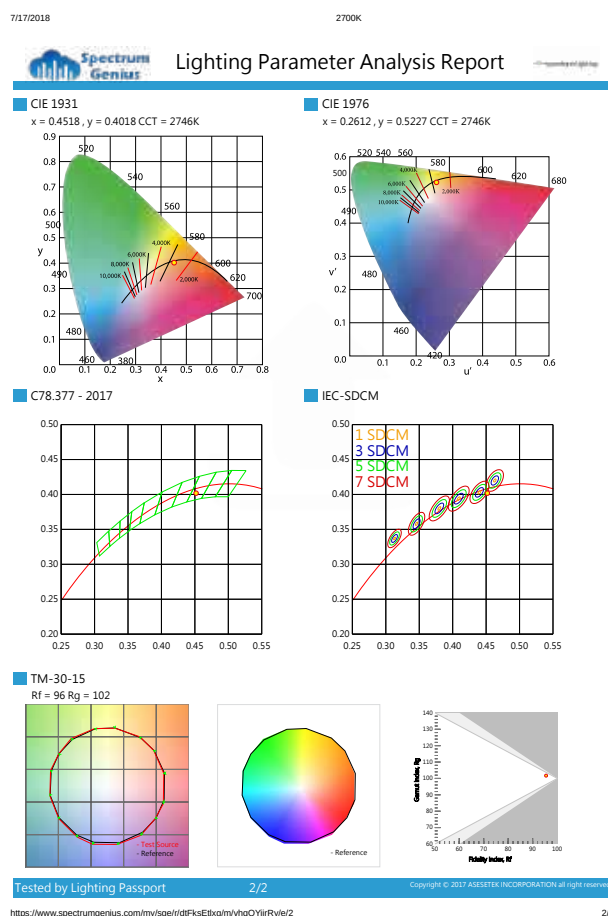
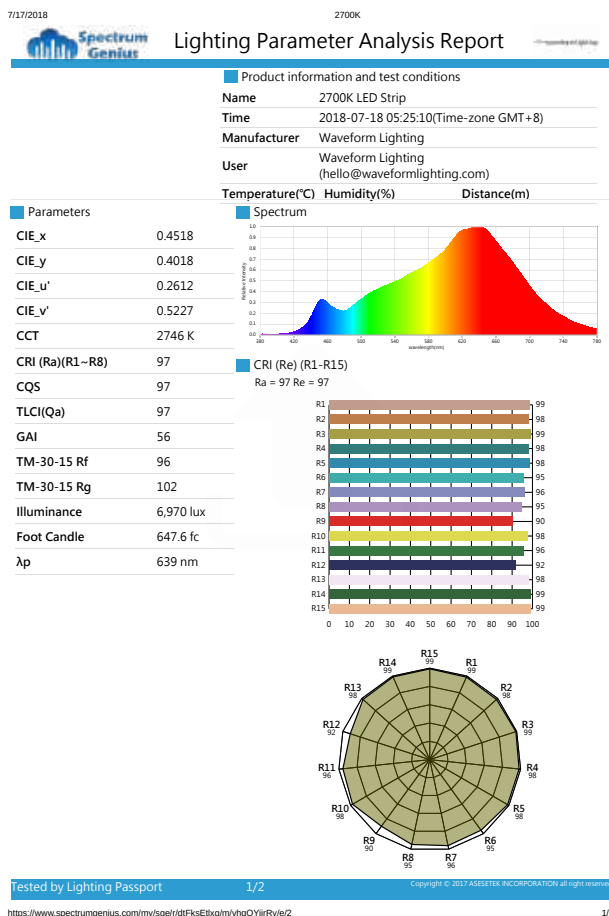
Na podstawie podanych powyżej informacji zainteresowani Czytelnicy mogą samodzielnie przeanalizować zdecydowaną większość danych, zaprezentowanych w przykładowym raporcie z testów taśmy LED o barwie białej cieplej (2700 K) – patrz **rysunek 14**.

Parametry rozrzutu chromatyczności

Warto dodać, że w profesjonalnych zastosowaniach pod uwagę brane są – oprócz opisanych wcześniej parametrów – także specjalne miary służące do oceny rozrzutu chromatyczności pomiędzy różnymi egzemplarzami źródeł światła, w tym. tzw. binning (wg normy C78.377) czy też tzw. elipsy MacAdama (SDCM – Standard Deviation Colour Matching). Zmniejszenie rozrzutu ma na celu ujednolicenie kolorystyki pomiędzy diodami pracującymi w ramach tego samego modułu czy też pomiędzy modułami (np. żarówkami zainstalowanymi we wspólnej oprawie czy w tym samym pomieszczeniu), ale do pewnego stopnia różnice parametrów spektralnych są nieuniknione. Dlatego też producenci stosują swego rodzaju „obejście” problemu: na etapie produkcji sortują bowiem diody w taki sposób, by do danego zbioru trafiały tylko egzemplarze różniące się nieznacznie pod względem chromatyczności. Często spotykana skala SDCM ma następującą postać:

- 1 SDCM – różnice barwy praktycznie pomijalne
- 2 SDCM – różnice barwy mierzalne, ale praktycznie niewidoczne
- 3 SDCM – różnice barwy słabo widoczne
- 4 SDCM – różnice barwy umiarkowanie widoczne
- 5 SDCM – różnice barwy wyraźnie widoczne

Na wykresie chromatyczności elipsy MacAdama obejmują obszary, w których gradient koloru jest praktycznie niezauważalny dla typowego oka ludzkiego – im wyższy numer elipsy, tym łatwiej dostrzegalne są różnice barwy (**rysunek 15**). Binning jest natomiast bezpośrednio powiązany ze znakowaniem poszczególnych partii diod LED – do każdego wycinka przestrzeni barw, leżącego w pobliżu krzywej ciała doskonale czarnego (w przypadku diod białych), przypisany jest określony kod alfanumeryczny, który pozwala na umiejscowienie chromatyczności tejże partii na wykresie CIE 1931 i ustalenie jej „odległości barwnej” od pozostałych grup (**rysunek 16**). Binning oraz skala SDCM to podstawowe narzędzia stosowane przez projektantów oświetlenia, pozwalające na dobór emiterów światła w taki sposób, by uzyskać jak najwyższą spójność chromatyczną – warto dodać, że wymogi w tym zakresie



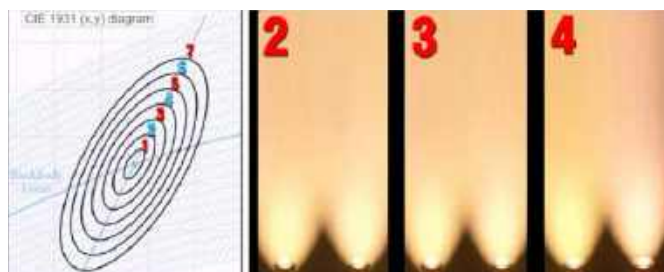
Rysunek 14. Przykładowy raport z testów białego paska LED (<http://t.ly/KOYIf>)

są wyższe dla oświetlenia pomieszczeń i znacznie mniej restrykcyjne dla oświetlenia zewnętrznego budynków.

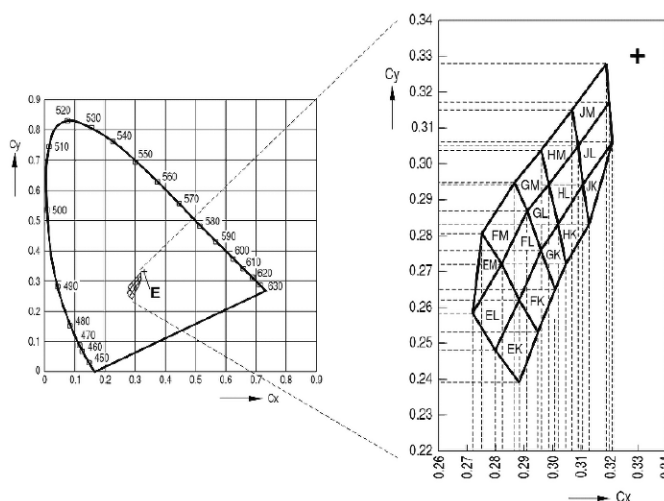
Parametry fotometryczne i radiometryczne źródeł LED

W opisach technicznych diod i modułów LED oraz gotowych urządzeń można spotkać się z szeregiem parametrów opisujących jasność oświetlenia, które często bywają mylone ze sobą. W oczywisty sposób prowadzi to do pewnego zamieszania, a w skrajnych przypadkach może być nawet wykorzystywane do manipulowania faktami – zwłaszcza przez dystrybutorów próbujących przekonać swoich odbiorców o niesamowitej wyższości oferowanych produktów nad wyrobami dostępnymi u konkurencji. Tymczasem dogłębne zrozumienie istoty parametrów fotometrycznych i różnic pomiędzy nimi a odpowiadającymi im wielkościami radiometrycznymi, pozwala znacznie pewniej poruszać się po zawiłościach optoelektroniki i samodzielnie weryfikować internetowe oferty.

Światłość [cd] to podstawowa wielkość fotometryczna ujęta w układzie SI i wyrażana w kandelach [cd]. Odnosi się do wizualnej jasności źródła (czyli natężenia strumienia) i – co ważne – jest ściśle powiązana ze standaryzowaną krzywą czułości ludzkiego oka. Innymi słowy – w przeciwieństwie do wielkości radiometrycznych (czyli określających możliwości źródła światła z energetycznego, czyli czysto fizycznego punktu widzenia) jest wielkością ważoną, zależną od długości fali. Odniesieniem pozwalającym na przeliczanie kandel na natężenie, wyrażone w watach na steradian [W/sr], jest czułość ludzkiego oka na falę o długości 555 nm (pik czułości). Nic więc dziwnego, że określanie kandel ma sens tylko w przypadku diod pracujących w świetle widzialnym, nie zaś dla emiterów IR czy UV. Ciekawostką jest fakt, że światłość 1 kandel odpowiada w przybliżeniu osiągom... „standardowej świeczki” o ściśle określonej masie i szybkości spalania, a wykonywanej



Rysunek 15. Przykład zastosowania skali SDCM do oceny różnic w barwie światła białego dwóch różnych lamp (<http://t.ly/TT0Zj>)



Rysunek 16. Przykładowy binning białych diod LED marki OSRAM (<http://t.ly/MTDmK>)

z wosku produkowanego na bazie spermacetu (substancji pozyskiwanej z głowy kaszalota). Choć może się to wydawać niewiarygodne, właśnie światło tak przygotowanej świecy było w dawnych latach wzorcem fotometrycznym, dopiero w 1948 roku zastąpionym kandelą ($1 \text{ świeca} = 1,02 \text{ cd}$).

Warto w tym miejscu zwrócić uwagę, że w przypadku nowoczesnych źródeł światła kandela jest używana przede wszystkim do określania mocy diod o charakterystyce silnie lub umiarkowanie kierunkowej (zwłaszcza małych emiterów w obudowach zintegrowanych z soczewką, w tym także dyfuzyjnych). W notach katalogowych niektórych diod z czołem płaskim można jednak znaleźć zarówno wartości światłości, jak i strumienia świetlnego – przyjrzyjmy się zatem drugiej z wymienionych wielkości.

Strumień świetlny [lm] to fotometryczna miara całkowitej ilości światła emitowanego przez źródło we wszystkich kierunkach. Miarą tej wielkości jest lumen [lm], przy czym strumień 1 lm jest wysyłany w jednostkowy kąt bryłowy (równy jednemu steradianowi, [sr]) przez izotropowe, punktowe źródło światła o światłości 1 cd, które umieszczone jest w wierzchołku tego kąta, czyli:

$$1 \text{ cd} = 1 \text{ lm/sr}$$

lub inaczej:

$$1 \text{ lm} = 1 \text{ cd} \cdot \text{sr}$$

Lumen jest zatem jednostką doskonale nadającą się do opisywania źródeł dookólnych (np. żarówek filamentowych), a także diod świecących w bardzo szerokim zakresie kątowym (z płaskim czołem) oraz popularnych modułów LED COB. Bardzo często strumień świetlny jest jednak (nie do końca sensownie) określany także dla urządzeń o silnie kierunkowej pracy, np. laterek. W praktyce, w celu przeliczenia lumenów na bardziej użyteczną jednostkę, trzeba byłoby określić także charakterystykę kierunkową urządzenia, wynikającą z zastosowanej optyki (a przynajmniej znać wartość kąta rozsyłu światła, np. dla połowy całkowitej mocy).

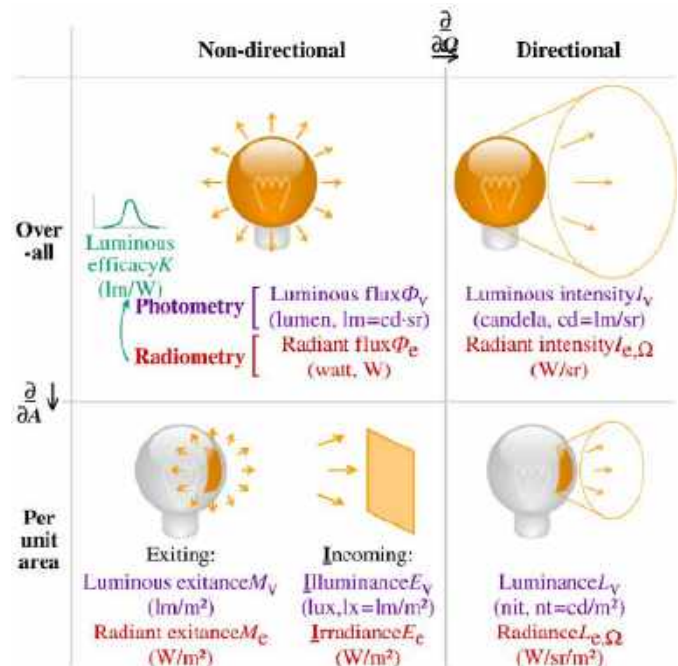
Natężenie światła [lx] to ostatni element podstawowej triady parametrów fotometrycznych, określający jasność, jaką dany strumień uzyskuje przy równomiernym rozłożeniu na dany obszar oświetlanej powierzchni.

$$1 \text{ lx} = 1 \text{ lm/m}^2$$

Podawanie wartości natężenia światła nie ma zatem sensu w przypadku opisu samego źródła – jest to bowiem wielkość zależna od odległości i charakterystyki kątowej całego układu optycznego: emiter(-y) + soczewki/dyfuzyory/odbłyśniki etc. Jeżeli jednak podamy nie tylko wartość natężenia, ale także odległość charakteryzowanej lampy od blatu roboczego, to podanie liczby luksów np. w środku pola roboczego będzie miało zdecydowany sens praktyczny, pozwoli bowiem odnieść się do obowiązujących norm BHP, które w bezpośredni sposób narzucają odpowiednie wymogi dotyczące oświetlenia stanowisk pracy (stosowne do różnych rodzajów wykonywanych czynności, np. pracy biurowej czy też precyzyjnych manipulacji).



Rysunek 17. Ewolucja skuteczności świetlnej diod LED na przestrzeni niecałej dekady (<http://t.ly/KcPk2>)



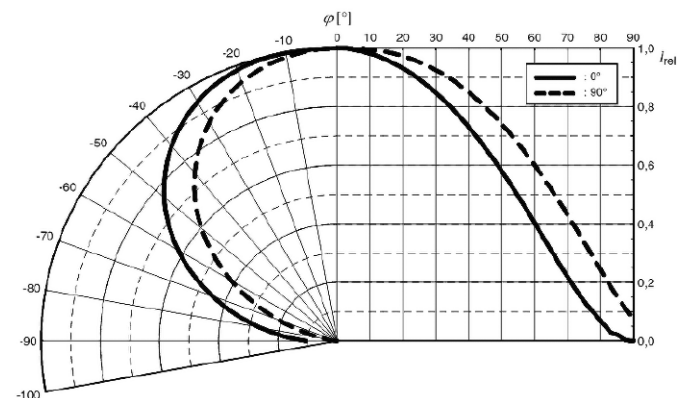
Rysunek 18. Porównanie najważniejszych parametrów fotometrycznych i radiometrycznych (<http://t.ly/mDSYA>)

Skuteczność świetlna [lm/W] jest wielkością w bardzo prosty sposób wiążącą opisaną powyżej parametry fotometryczne z „twardymi faktami”, czyli radiometrią. Liczba ta określa stosunek strumienia świetlnego, emitowanego przez dane źródło światła, do mocy zasilania tegoż źródła. Co interesujące, skuteczność świetlna może być podawana w odniesieniu do tzw. źródła idealnego, czyli hipotetycznego źródła o długości fali równej 555 nm i o sprawności równej 100%. W przypadku takiego teoretycznego tworu mielibyśmy do czynienia z wartością równą 683 lm/W. Jako ciekawostkę warto podać, że w 2014 firma Cree – producent popularnych diod dużej mocy – osiągnęła rekord świata, produkując nową generację białych emiterów LED o skuteczności 303 lm/W (rysunek 17). Jak widać wciąż „trzymamy się” daleko poniżej połowy teoretycznej wydajności maksymalnej.

Doskonale porównanie wielkości fotometrycznych i radiometrycznych, określanych dla źródeł izotropowych oraz kierunkowych, pokazano na rysunku 18. Warto znać i rozumieć zaprezentowane na infografice zależności pomiędzy poszczególnymi jednostkami i wielkościami.

Charakterystyka kątowa

W dotychczasowym opisie jedynie delikatnie zaznaczyliśmy kwestię charakterystyk kątowych diod LED. Tymczasem w praktyce nierzadko okazują się one równie ważne (a w pewnych przypadkach



Rysunek 19. Charakterystyka kątowa diody KW DELPS2.RA marki Osram (<http://t.ly/484Ur>)

nawet ważniejsze) od parametrów spektralnych czy fotometrycznych. Na **rysunku 19** pokazano przykładowe charakterystyki kątowe niewielkiej diody LED SMD, wyrażone postaci podwójnych wykresów. Lewa strona to klasyczny wykres w biegunowym układzie współrzędnych, zaś prawa – we współrzędnych kartezjańskich. Obydwie części rysunku prezentują dokładnie te same dane



Fotografia 13. Dioda LED z serii TOPLED® E1608 marki OSRAM (<http://t.ly/VIS3S>)

– pozioma siatka prawej strony jest bowiem stycznym przedłużeniem siatki polarnej. W tym przypadku producent wykreślił dwie krzywe, odpowiadające ortogonalnym względem siebie płaszczyznom, prostopadłym do powierzchni płytki drukowanej. Różnice są nieprzypadkowe – grafika pochodzi bowiem z noty katalogowej diody prostokątnej typu TOPLED® E1608 marki OSRAM (**fotografia 13**). Pionowa oś (y) reprezentuje względne natężenie promieniowania (w odniesieniu do natężenia zmierzonego dokładnie na osi optycznej). Warto dodać, że liczbowa wartość kąta emisji światła danej diody, którą można znaleźć na początku każdej noty katalogowej, jest zwykle określana jako „całkowity” kąt rozsyłu w danej płaszczyźnie (w tym przypadku 110° dla płaszczyzny 0° i 130° dla płaszczyzny 90°), albo jako kąty połowiczne (odpowiednio: $\pm 55^\circ$ i $\pm 65^\circ$). Podane wartości są wyznaczane z punktów przecięcia

odpowiednich krzywych z poziomą linią odpowiadającą $I_{rel}=50\%$. Kąt całkowity nosi nazwę szerokości w połowie wysokości, ang. *full width at half maximum* (FWHM), zaś w drugim opisanym przypadku można spotkać się z nazwą ang. *angle of half sensitivity*, czyli „kąt połowicznej czułości”. Co ważne, w katalogach producentów wartości te są często (i niestety błędnie) stosowane zamiennie. Przykładowo: dioda o kącie połowicznej czułości $\pm 60^\circ$ może być wrzucona „do jednego worka” z emiterami, dla których 60° jest w istocie kątem FWHM – to kolejny przykład na to, że automatycznym filtrem dostępnym w bazach dystrybutorów komponentów nie należy zbyt ufać i zawsze warto sprawdzić wszystkie istotne parametry w nocie katalogowej wybranego komponentu...

Podsumowanie

W artykule zaprezentowaliśmy kilkanaście aspektów związanych z charakterystykami optycznymi diod LED i modułów LED COB. Większość opisanych zagadnień w równym stopniu dotyczy także gotowych opraw ze zintegrowanym źródłem LED, jak i innych urządzeń – w tym latarek, oświetlenia rowerowego, lamp fotograficznych, a nawet niektórych urządzeń oraz akcesoriów medycznych. Mamy nadzieję, że niniejszy poradnik pozwoli naszym Czytelnikom znacznie sprawniej poruszać się przez gąszcz parametrów foto- i radiometrycznych, a także rozmaitych charakterystyk dodatkowych – tym bardziej, że noty katalogowe współczesnych diod LED osiągają dziś długość kilkudziesięciu stron, z czego zdecydowana większość to właśnie opisy dotyczące samej tylko optyki.

inż. Przemysław Musz, EP

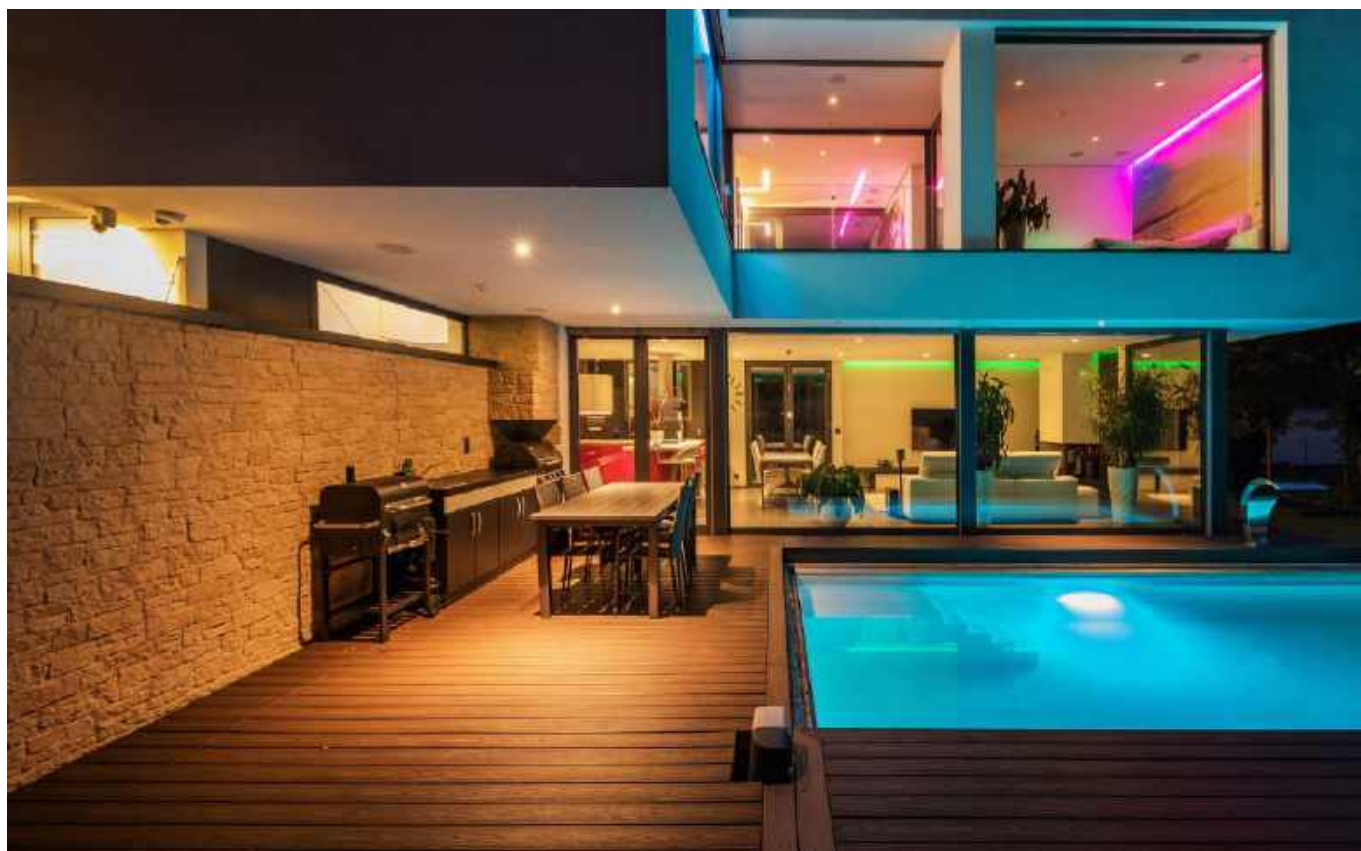
REKLAMA







przejrzyj i kup na
www.ulubionykiosk.pl



Oświetlenie LED w aplikacjach konsumenckich i profesjonalnych

Diody LED zrewolucjonizowały i całkowicie zdominowały branżę oświetleniową. Te energooszczędne, trwałe, bezpieczne i niezwykle elastyczne pod względem aplikacyjnym komponenty rozświetlają nasze miasta – fasady oraz wnętrza domów mieszkalnych i biurowców, szerokie arterie i małe uliczki, mosty, instalacje artystyczne, zabytki... a to zaledwie początek niekończącej się listy zastosowań źródeł elektroluminescencyjnych. Diody LED szturmem wdarły się do branży motoryzacyjnej, medycznej czy nawet... rolniczej. W artykule bierzemy „na warsztat” wybrane przykłady zastosowania najnowocześniejszych diod LED oraz specjalistycznych technik, mających na celu niemal idealne dostrojenie parametrów spektralnych wynikowego strumienia świetlnego. Umyślnie pomijamy przy tym zastosowania podstawowe, np. zwykłe żarówki i taśmy LED, dostępne w każdym markecie AGD/RTV – skupiamy się na najciekawszych naszym zdaniem przykładach współczesnych systemów oświetlenia.

Oświetlenie pomieszczeń

Wraz z upowszechnianiem żarówek, taśm czy też modułów opartych na białych diodach LED, sukcesywnie rośnie znaczenie optymalizacji spektrum oświetlenia. Użytkownicy – również prywatni – oczekują efektów możliwie jak najbardziej zbliżonych do naturalnego światła słonecznego, co doskonale wpisuje się w nurt określany mianem *wellbeing*. Co ciekawe: w literaturze (także polskiej), poświęconej zagadnieniom BHP, można odnaleźć interesujące rozważania na temat potencjalnych zysków biologicznych, jakie możemy osiągnąć dostosowując barwę światła w pomieszczeniach podczas pracy zmianowej. Kwestia oświetlenia „podążającego” za naturalnym zegarem biologicznym i zmieniającego się zgodnie z rytmem okołodobowym (cyrkadialnym) jest zatem podnoszona już nie tylko w odniesieniu do oświetlenia wnętrz domowych.

Nasz ekspresowy przegląd rozpoczniemy od jednego z najprostszych rozwiązań pozwalających na uzyskanie światła o szerokim



Fotografia 1. Przykładowa taśma LED złożona z diod białych o barwie ciepłej i zimnej (<http://t.ly/hto1P>)

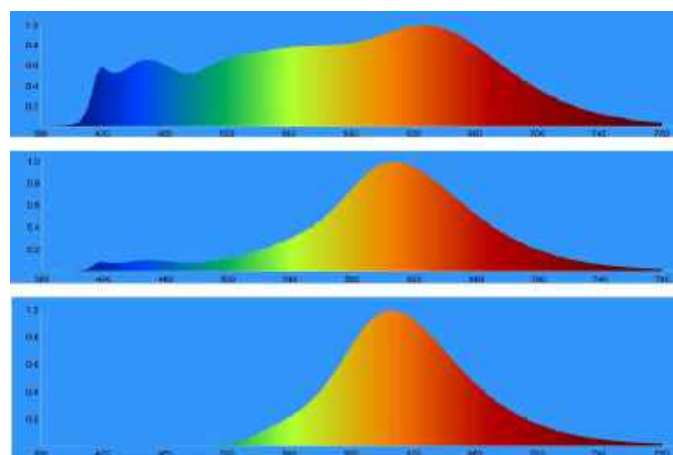
zakresie temperatury barwowej. Niektóre taśmy LED – takie jak zaprezentowana na **fotografii 1** (model StudioFlex CT Adjustable LED Strip) są złożone z dwóch rodzajów diod, różniących się temperaturą barwową. W zależności od stosunku wysterowania obydwu sekcji można uzyskiwać CCT na poziomie od 2700 K do 6500 K, przy CRI przekraczającym 90.

Jeszcze wyższą jakość widma, znacznie bardziej zbliżoną do naturalnego światła słonecznego, można uzyskać przy zastosowaniu specjalistycznych żarówek lub opraw LED, określanych zwykle mianem *full spectrum*, *sunlight-like* itp. Niektóre z nich są przestrajalne i umożliwiają przełączanie pomiędzy trzema predefiniowanymi kształtami widma optycznego – przykładem może być tutaj zaawansowana oprawa BioLight™ Recessed Light, pokazana na **fotografii 2**. W zależności od ustawienia pozwala ona uzyskać barwę ciepłą, zimną lub neutralną – ta ostatnia z doskonałą dokładnością replikuje widmo światła słonecznego (**rysunek 1**).

Znakomite osiągi (choć już przy jednym, ustalonym kształcie widma) oferują zaawansowane żarówki marki Yuji Lighting,



Fotografia 2. Oprawa LED Bio-Light™ Recessed Light (<http://t.ly/kRz63>)



Rysunek 1. Porównanie widm optycznych uzyskiwanych przy zastosowaniu oprawy LED BioLight™ Recessed Light w trybach pracy: Full Spectrum Day Mode (na górze), Low Blue Mixed Mode (w środku), No Blue Night Mode (na dole). Źródło: <http://t.ly/kRz63>

sprzedawane pod nazwą SunWave PAR30 Dimmable (**fotografia 3**). Oprócz doskonałego odwzorowania kolorów (CRI ≥ 98 , R9 ≥ 95) produkt umożliwia ściemnianie nawet do 5% maksymalnej jasności, zaś za chłodzenie (pozytywnie wpływające na termiczną stabilność widma) odpowiada masywny, wbudowany radiator.

Oświetlenie do upraw roślin

Konieczność dostarczenia roślinom odpowiedniej ilości światła jest tak oczywista, jak wymóg właściwego dla danego gatunku nawodnienia. Oświetlenie LED o specyficznie ukształtowanym widmie umożliwia dodatkowe stymulowanie wzrostu roślin poprzez zwiększenie efektywności procesów biologicznych, w tym przede wszystkim fotosyntezy i tzw. fotomorfogenezy. Najważniejsze pasma to przede wszystkim światło niebieskie (400...500 nm) i czerwone (600...700 nm).

Światło niebieskie (głównie w zakresie ok. 450 nm) wspomaga rozwój liści i jest istotne na wczesnych etapach wzrostu, stymuluje bowiem fotosyntezę oraz wzrost wegetatywny. Z kolei światło czerwone (ok. 660 nm) odpowiada za procesy związane z kwitnieniem i owocowaniem. Dodatkowo często stosuje się światło leżące na granicy pasma widzialnego oraz bliskiej podczerwieni (nawet do 750 nm), które wpływa na wydłużanie pędów i procesy regulacji fotoperiodycznej. Przedmiotem zainteresowania naukowców pozostaje także wpływ ultrafioletu na rozwój roślin – okazuje się bowiem, że oddziałuje on na aromat i kolorystykę kwiatów, a niektóre źródła podają nawet korzystny wpływ UV na produkcję niektórych związków chemicznych o znaczeniu leczniczym.

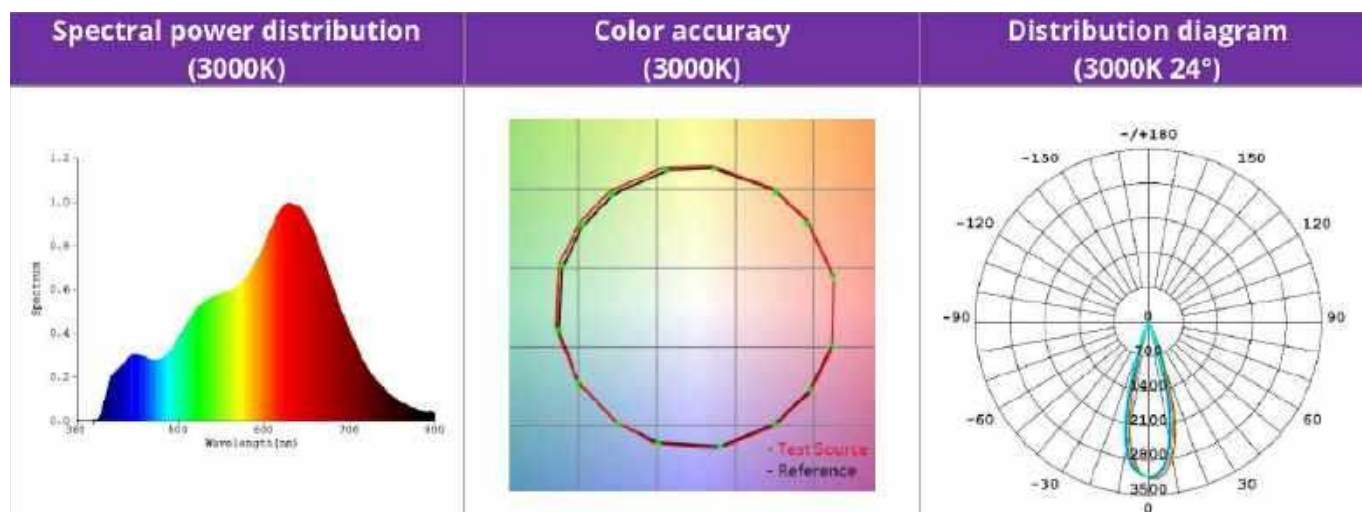
Najprostsze (ale zarazem bardzo skuteczne) spośród dostępnych na rynku rozwiązań z tej kategorii opierają się na zastosowaniu specjalnych żarówek



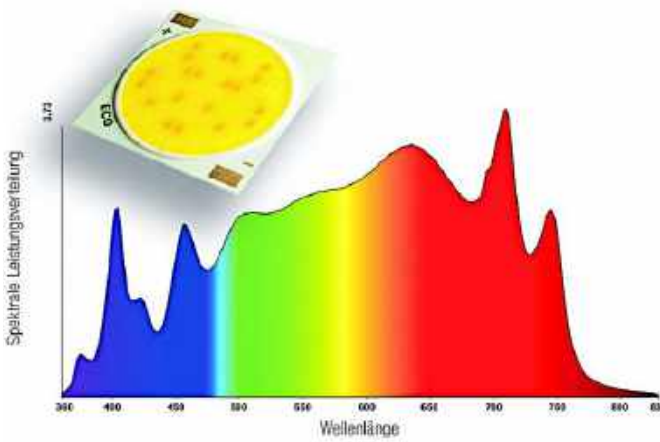
Fotografia 3. Żarówka z serii SunWave PAR 30 Dimmable marki Yuji Lighting (<http://t.ly/gwo28>)



Fotografia 4. Żarówka LED przeznaczona do stymulacji wzrostu roślin (<http://t.ly/qduix>)



Rysunek 2. Parametry spektralne żarówki z serii SunWave PAR 30 Dimmable o temperaturze barwowej 3000 K (<http://t.ly/gwo28>)



Rysunek 3. Moduł LED COB typu GW1919U niemieckiej marki euroLighting (<http://t.ly/OU9S9>)

złożonych z zestawu czerwonych i niebieskich diod LED (fotografia 4). „Wycięcie” fragmentu widma obejmującego m.in. światło zielone ma przy okazji spore znaczenie z energetycznego punktu widzenia – promieniowanie w tym zakresie jest wszak praktycznie w całości odbijane przez chlorofil typu A (zawarty w komórkach



Fotografia 5. Widok modułu oświetleniowego marki Mammoth Lighting (<http://t.ly/iwqtg>)



Fotografia 6. Kompletny system oświetlenia upraw roślinnych marki Mammoth Lighting (<http://t.ly/iwqtg>)

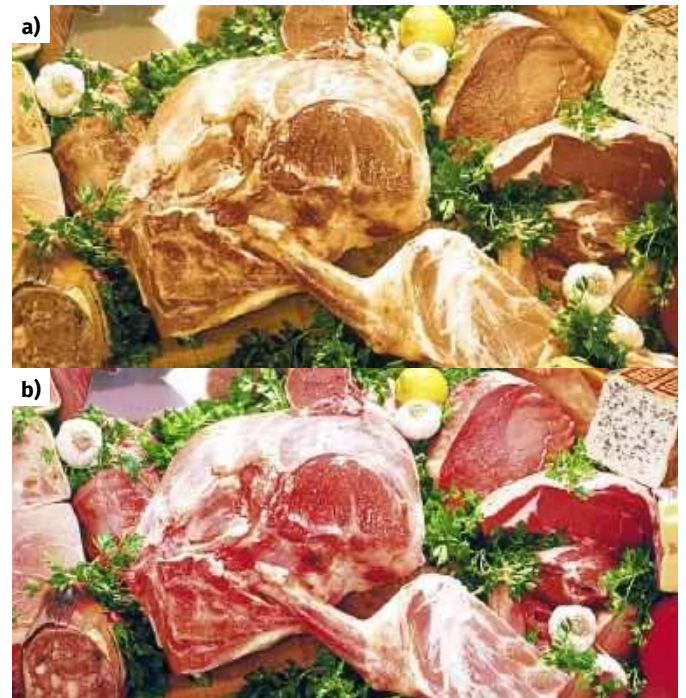
roślinnych), a zatem długotrwałe naświetlanie nim roślin miałyby w podstawowych zastosowaniach mocno ograniczony sens praktyczny.

Bardziej rozbudowane systemy mogą opierać się na specjalistycznych emiterach lub modułach LED COB – jako szczególnie ciekawy przykład przywołamy tutaj moduł GW1919U niemieckiej marki euroLighting (rysunek 3), który – oprócz odpowiednio zoptymalizowanego widma widzialnego – emituje także promieniowanie w zakresie ultrafioletu (UV A). W najbardziej zaawansowanych aplikacjach zastosowanie znajdują już rozbudowane systemy oparte na idei oświetlenia hybrydowego, łączącego w sobie zestaw diod białych i monochromatycznych (fotografia 5), zabudowanych w podłużnych modułach i sterowanych przez zintegrowane kontrolery (fotografia 6).

Oświetlenie ekspozycyjne dla gastronomii i handlu spożywczego

Jeżeli zapytalibyśmy klientów nowoczesnego marketu spożywczego, ekskluzywnej piekarni czy też dobrego sklepu mięsnego o to, czy coś zaciekało ich w sposobie oświetlenia ekspozycyjnych produktów, to zapewne tylko nieliczni (może 1 osoba na 1000?) zwróciliby uwagę na zastosowane tam lampy. Tymczasem właściciele sklepów doskonale wiedzą, że klienci bardzo często „kupują oczami”, a jakość oświetlenia – mierzona zarówno intensywnością, jak i charakterystyką spektralną – ma ogromne znaczenie dla wizualnej atrakcyjności oferowanych produktów. Praktycznie każdy rodzaj towaru wymaga nieco innego oświetlenia, dlatego skupimy się tylko na wybranych przykładach.

1. **Oświetlenie mięsa i wędlin** powinno zawierać zwiększony udział składowych w paśmie bliskiej i dalekiej czerwieni, choć (pomijając wspomniane podbicie w zakresie najdłuższych fal) podstawowa część widma odpowiada zwykle chłodnemu światłu białemu. W efekcie wynikowy strumień sprawia wrażenie biało-różowego (fotografia 7). Co ciekawe, niektóre badania wykazują związek odpowiednio dobranej barwy oświetlenia lądowej ekspozycyjnej z szybkością psucia się surowego mięsa. Bardzo istotne jest także ograniczenie emisji ciepła, co pozwala



Fotografia 7. Ilustracja wpływu barwy oświetlenia na wizualną atrakcyjność mięsa ekspozowanego w gablocie (a – światło ciepłe, b – światło chłodne z podbiciem pasma w rejonie czerwieni – materiały marki Promolux (<http://t.ly/Y1bjl>))

utrzymać właściwe warunki termiczne w zamkniętej i stosunkowo niewielkiej przestrzeni pod szybą lacy.

2. Oświetlenie pieczywa i serów daje najlepsze efekty wizualne w przypadku zastosowania źródeł o bardzo niskiej temperaturze barwowej (poniżej 2700 K), podkreślającej odcienie żółci, brązu i czerwieni. Ze względu na sposób prezentacji towaru w praktyce stosowane są zwykle odpowiednio ustawione reflektory (fotografia 8) lub listwy LED podświetlające poszczególne półki regałów ekspozycyjnych. Ciepła barwa światła potęguje wizualnie efekt świeżości i skutecznie zachęca klientów do dokonywania zakupów (fotografia 9).

3. Oświetlenie lodów i wyrobów cukierniczych, ze względu na duże zróżnicowanie barw, należy dobierać w taki sposób, by dobrze odwzorowywało kolorystykę zastosowanych w produktach barwników (fotografia 10). Duże znaczenie ma tutaj



Fotografia 8. 1-fazowy reflektor szynowy Greenie 2200 K (http://t.ly/_A0ej)



Fotografia 9. Przykładowa ekspozycja piekarni z zastosowaniem reflektorów o temperaturze barwowej 2200 K (http://t.ly/_A0ej)



Fotografia 10. Odpowiednie odwzorowanie barw jest konieczne w przypadku ekspozycji wyrobów cukierniczych oraz lodów (<http://t.ly/4u2we>)

odpowiednio wysoki wskaźnik CRI, choć istnieją także pewne szczególne wymagania dotyczące części widma widzialnego oraz promieniowania UV – okazuje się, że niektóre związki zawarte w słodyczach (np. aminokwasy czy tłuszcze) mogą wchodzić w niepożądane reakcje chemiczne pod wpływem określonych długości fali. Dotyczy to także barwników, które – zwłaszcza w wyniku działania ultrafioletu – ulegają odbarwieniu po odpowiednio długiej ekspozycji.

Oświetlenie LED do aplikacji profesjonalnych

Omówiliśmy do tej pory wybrane aplikacje konsumenckie i komercyjne, które doskonale ilustrują zagadnienia opisane w artykule przewodnim tego wydania „Elektroniki Praktycznej” pt. *Parametry optyczne źródeł światła LED*. Teraz przechodzimy do zastosowań bardziej zaawansowanych, wymagających najwyższej jakości źródeł światła o ściśle kontrolowanych charakterystykach widmowych.

Fotografia i kinematografia

Współczesna fotografia i kinematografia coraz silniej przestawiają się na użycie półprzewodnikowych źródeł światła, odchodząc jednocześnie od rozwiązań klasycznych, np. wyładowczych lamp błyskowych czy żarówek dużej mocy (w przypadkach wymagających zastosowania światła ciągłego, np. na planach filmowych czy w trakcie sesji zdjęciowych w stylu fashion). Wierne odwzorowanie naturalnego koloru skóry (ale także tła czy rekwizytów) jest niezwykle istotne z prostej przyczyny: niewłaściwie dobrana charakterystyka spektralna niechybnie spowoduje „obcięcie” pewnych przebiegów tonalnych, których odzyskanie na etapie postprocessingu (retuszu zdjęć czy montażu filmu) może być często niewykonalne. Dlatego też profesjonalni fotograficy i oświetleniowcy coraz częściej korzystają z nowoczesnych lamp LED, których zaletą – oprócz dobrych charakterystyk widmowych – jest także znacznie mniejsza awaryjność (większa żywotność) oraz... niższa emisja ciepła podczas pracy ciągłej. Ta ostatnia zaleta ma szczególne znaczenie w przypadku stosunkowo ciasnych przestrzeni (np. niewielkich studiów fotograficznych), w których warunki – w wyniku oświetlenia za pomocą konwencjonalnych lamp wielkiej mocy – po pewnym



Fotografia 11. Przykładowy reflektor kinematograficzny z segmentu niskobudżetowego (<http://t.ly/GPwYh>)

czasie stałyby się nie do wytrzymania dla fotografowanych modelek/-i, ale także (po części) dla samej ekipy.

Warto zwrócić uwagę, że wysokiej jakości oświetlenie LED mieści się w budżecie już zamożniejszych amatorów oraz początkujących fotografów zawodowych. Kompaktowe, lekkie, a zarazem mocne reflektory są dostępne w cenach oscylujących wokół 1000...2500 zł, a co ważne – niektóre modele współpracują z szeroką gamą akcesoriów pomocniczych – statywami, systemami zdalnego wyzwalania, kompatybilnymi czasami, a także wszelkiego rodzaju modyfikatorami światła: parasolkami czy softbokami. Niektóre z nich mają możliwość regulacji temperatury barwowej w szerokim zakresie (**fotografia 11**), inne pracują z ustaloną wartością CCT, zwykle w rejonie odcieni neutralnych (np. 5600 K – **fotografia 12**).

Oświetlenie LED znalazło także liczne zastosowania w aplikacjach telewizyjnych, w których ogromna liczba źródeł półprzewodnikowych (rozieszczonych w różnych punktach studia) pozwala zmniejszyć koszty serwisowe i uprościć obsługę logistyczną. W tego typu aplikacjach stosowane są zarówno reflektory, jak i panele LED, zapewniające bardziej rozproszone, miękkie i naturalne światło (**fotografia 13**). Użycie dodatkowych modyfikatorów w postaci kratownic lub tzw. plastrów miodu umożliwia dodatkowe ukierunkowanie



Fotografia 12. Lampa studyjna LED Godox SL-200 W II (<http://t.ly/LSRtW>)



Fotografia 15. Profesjonalny reflektor studyjny Aputure LS 600x Pro (<http://t.ly/GsCMM>)

światła i – w dalszej kolejności – lepszą kontrolę nad poszczególnymi planami (**fotografia 14**).

Do źródeł LED powoli przekonują się także filmowcy czy profesjonalni fotograficy. Najnowocześniejsze lampy pozwalają nie tylko na regulację temperatury barwowej (przy bardzo wysokim CRI oraz TLCI (tzw. Television Lighting Consistency Index), ale także na bezstopniowe sterowanie jasnością w krokach nawet co 0,1% – takie parametry oferuje przykładowy reflektor Aputure LS 600x Pro (**fotografia 15**), umożliwiając uzyskanie natężenia światła równego nawet 18510 lx (!) na odległości 3 m (przy zastosowaniu specjalnej soczewki Fresnela).

Zastosowania medyczne

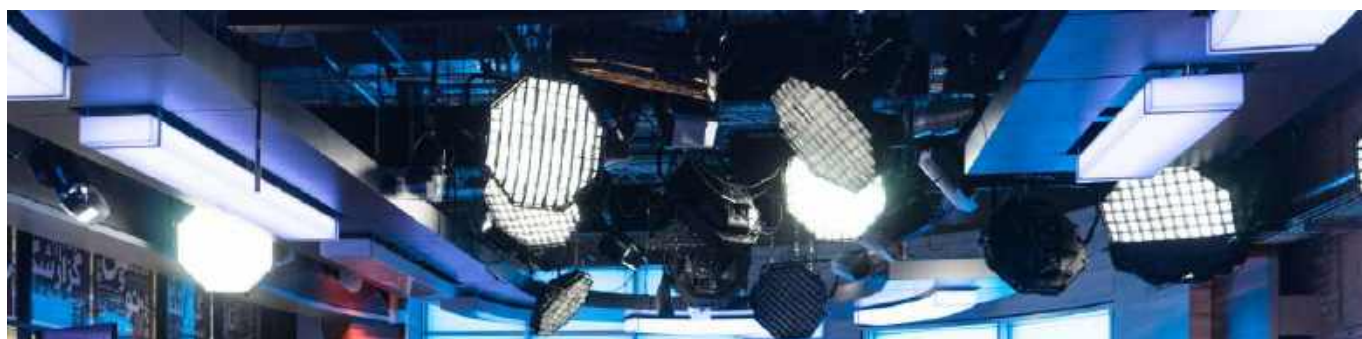
Na koniec naszego przeglądu pozostawiliśmy aplikacje medyczne – tę część prezentacji rozpoczniemy od niepozornych, choć w istocie niezwykle ważnych przyrządów, jakimi są niewielkie oświetlacze przeznaczone do celów diagnostycznych.

Wyobraźmy sobie, że mamy do zaprojektowania otoskop. Przyrząd relatywnie prosty i – jak na medyczne standardy – raczej niedrogi. Jak dobrać do niego źródło światła? Nie ulega wątpliwości, że diody LED będą najbardziej odpowiednim rozwiązaniem – niewielki pobór mocy, bardzo ograniczone wydzielanie ciepła, kompaktowe wymiary... same zalety. Należy jednak wziąć pod uwagę, że tego typu przyrządy są stosowane do wzrokowej oceny stanu tkanek ucha środkowego oraz całego przewodu słuchowego – wszelkie stany zapalne czy zmiany patologiczne muszą zatem być dokładnie odwzorowane, gdyż przejścia tonalne mają znaczenie diagnostyczne: pozwalają określić m.in. granice obszaru zajętego chorobowo, a także rodzaj czy stopień zaawansowania obecnych w uchu zmian. Dlatego wierność oddawania barw jest tutaj niebywale istotna – dotyczy to zwłaszcza czerwieni i koloru cielistego, czyli w grę wchodzi przynajmniej odpowiednio wysokie CRI oraz indeks R9. Bardzo ważna jest także jednorodność oświetlenia – niedopuszczalne jest, by np. na granicach obrazowanego pola w widoczny sposób zmieniała się temperatura barwowa strumienia.

Jako przykład dostępnego na rynku urządzenia tego typu przywołajmy markowy otoskop światłowodowy BETA 200 LED firmy Heine (**fotografia 16**). Zastosowany w nim oświetlacz LED emituje jednorodną wiązkę światła o temperaturze barwowej 3500 K, CRI > 97 i R9 > 93. Natężenie uzyskiwanego dzięki wbudowanej optyce światła wynosi 77000 lx, zaś jasność może być regulowana w zakresie od 3% do 100%.



Fotografia 13. Zastosowanie lamp LED do oświetlenia studia telewizyjnego (<http://t.ly/uMe-a>)



Fotografia 14. Panele LED z kratownicami (<http://t.ly/Lf944>)

Analogiczny przykład zastosowania nowoczesnego źródła LED stanowi inny produkt, także ze stajni firmy Heine – laryngologiczna lampa czołowa ML4 LED (**fotografia 17**). W tym przypadku wbudowany oświetlacz musi być rzecz jasna nieporównanie silniejszy, niż w otoskopie, gdyż obrazowane tkanki znajdują się w pewnej odległości (rzędu ~20 cm) od twarzy lekarza. Oświetlacz modelu ML4 LED oferuje oświetlenie na poziomie 90000 lx przy odległości 180 mm, a dodatkowo wbudowany układ optyczny pozwala na płynną regulację średnicy pola oświetlenia w zakresie od 30 do 80 mm, przez co urządzenie jest przydatne także dla lekarzy z branż innych niż laryngologia (m.in. dentystów).

Idąc dalej – w kierunku coraz bardziej zaawansowanych urządzeń medycznych – trafiamy na mikroskopy operacyjne, będące



Fotografia 16. Otoskop marki Heine z oświetlaczem LED (<http://t.ly/ZMJ5o>)



podstawowym wyposażeniem nie tylko chirurgów (neurochirurgia, chirurgia okulistyczna i otolaryngologiczna, mikrochirurgia rekonstrukcyjna i replantacyjna, chirurgia plastyczna itp.), ale także niektórych stomatologów. Najbardziej rozbudowane mikroskopy, stosowane w wymagających procedurach zabiegowych, są dziś wyposażane w wysokiej klasy oświetlacze LED, często oferujące funkcje przeznaczone do użycia w ściśle określonych branżach. Przykład? Ultranowoczesny mikroskop Proveo 8 marki Leica (**fotografia 18**) – jednego ze światowych potentatów na rynku przyrządów optycznych i fotograficznych – jest wyposażony w specjalistyczny moduł LED o wysokiej wydajności, którego zadanie wykracza daleko poza zwykłe oświetlenie pola roboczego. Producent zastosował bowiem innowacyjny system koncentrycznego, cztero-wiązkowego oświetlenia z regulowanym zoomem, umożliwiający stabilne wytwarzanie efektu czerwonego odbicia od struktur soczewki oka w czasie operacji katarakty – ów efekt ma fundamentalne znaczenie dla operatora, bowiem to właśnie ten niepozorny artefakt pozwala zobrazować wszelkie, najmniejsze nawet detale usuwanej soczewki i znacznie ułatwia bieżącą ocenę efektów zabiegu podczas implantacji soczewki sztucznej.

I jeszcze gratka dla wszystkich dociekliwych Czytelników, którzy lubią „wiedzieć, co w trawie piszczy”. Konstrukcja opisywanego modułu oświetlacza okazała się dla inżynierów z firmy Leica nie lada wyzwaniem i to – co ciekawe – nie ze względu na właściwości spektralne, ale na... pobór mocy! Ze względów praktycznych konieczna była bowiem silna miniaturyzacja całości, co przy zastosowaniu umiejscowionych tuż koło siebie, dwóch diod LED dużej mocy oraz trzeciego (znajdującego się pomiędzy nimi) modułu COB, było bardzo problematyczne. Sam emiter COB pracuje bowiem z mocą 8 W, a cały moduł pobiera podczas pracy aż 27 W – przy bardzo kompaktowych wymiarach całości (**fotografia 19**) niezbędne okazało się zastosowanie bloku chłodzącego (**fotografia 20**). Problem okazał się



Fotografia 17. Profesjonalna lampa czołowa marki Heine – model ML4 LED (<http://t.ly/J06g9>)



Fotografia 18. Nowoczesny mikroskop operacyjny Proveo 8 marki Leica (<http://t.ly/3dBPR>)



Fotografia 19. Moduł oświetlenia zastosowany w mikroskopie chirurgicznym Proveo 8 marki Leica (<http://t.ly/lcAEw>)



tak trudny do rozwiązania, że producent musiał uciekać się do zaawansowanego oprogramowania do symulacji termicznych, opracowanego przez naukowców z Uniwersytetu Technicznego w Wiedniu.

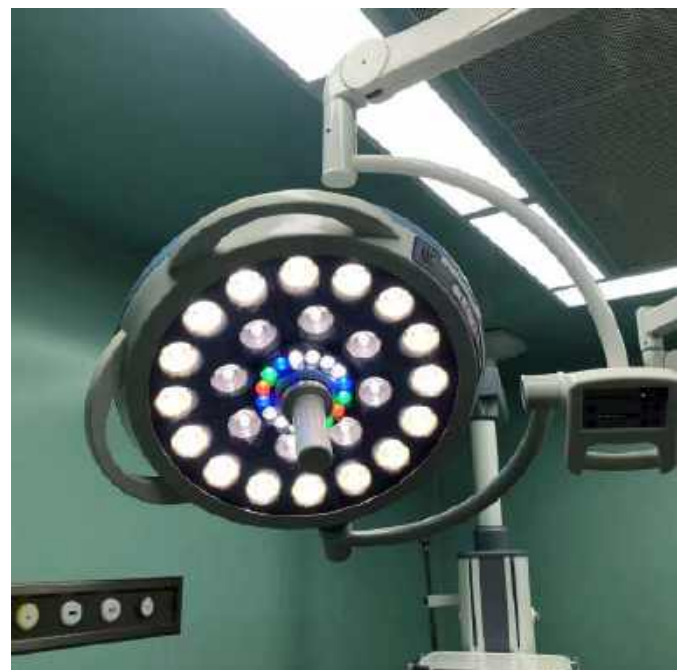
Zostawmy na chwilę operacje przeprowadzane w ogromnych powiększeniach – zabiegi chirurgii otwartej także wymagają bowiem perfekcyjnego (pod względem widma i rozkładu przestrzennego) oraz bardzo wydajnego oświetlenia, które w nowoczesnych szpitalach jest już coraz częściej realizowane przez źródła LED-owe. Nowoczesne lampy operacyjne są spektakularnym przykładem „żonglowania” parametrami oświetlenia hybrydowego. Prawidłowe odwzorowanie barw pozostaje niebawie istotne dla chirurgów wykonujących zabiegi otwarte, pozwala bowiem nie tylko łatwiej różnicować poszczególne struktury anatomiczne, ale także umożliwia bardziej wiarygodne odnajdywanie granic tkanek zmienionych patologicznie (np. guzów nowotworowych) – dla doświadczonego chirurga ogromne znaczenie mają bowiem nie tylko wrażenia dotykowe (palpacja tkanek, opór stawiany narzędziom chirurgicznym), ale także wzrokowe. Dlatego też oświetlenie stosowane na salach operacyjnych musi być nie tylko bardzo intensywne, ale zwykle także bezcieniowe (dzięki odpowiedniemu ułożeniu poszczególnych modułów), zaś stopień odwzorowania barw (zwłaszcza w zakresie czerwieni) musi być „podkręcony” do maksimum. Na **fotografii 21** można zobaczyć przykładową lampę operacyjną wykonaną w technologii hybrydowej – łączącej przemysłowy zestaw źródeł dużej mocy (rozieszczonych na planie dwóch koncentrycznych

okręgów) oraz dodatkowych diod RGB znajdujących się wokół centralnego uchwytu i służących do modyfikowania temperatury barwowej w zależności od aktualnych potrzeb operatora. Urządzenie umożliwia uzyskanie natężenia oświetlenia na poziomie od 60000 do nawet 160 000 lx, zaś temperatura barwowa może być regulowana w zakresie od 3500 do 5000 K (przy CRI > 95). Co ważne, tak silny strumień światła powoduje tylko nieznaczne podniesienie temperatury otoczenia (nad głową chirurga) oraz pola operacyjnego – producent podaje wartości, odpowiednio, 1°C i 2°C, co ma niebagatelne znaczenie zwłaszcza podczas wielogodzinnych operacji (w ramach jednego z projektów autor sam miał okazję uczestniczyć w zabiegu trwającym 12 godzin, a wcale nie jest to jeszcze rekord w tym obszarze medycyny zabiegowej!).

Z kolei **fotografia 22** prezentuje bezcieniową lampę z serii ALYON marki Steris, wyposażoną w 14 modułów złożonych z naprzemiennie ułożonych emiterów o różnych temperaturach



Fotografia 20. Montaż bloku chłodzącego modułu z fotografii 19 (<http://t.ly/lcAEw>)



Fotografia 21. Nowoczesna lampa operacyjna LED (<http://t.ly/NN5F2>)



Fotografia 22. Inteligentna lampa operacyjna ALYON marki Steris (<http://t.ly/xYXnB>)



Fotografia 23. Inteligentna lampa operacyjna ALYON marki Steris (<http://t.ly/xYXnB>)

barwowych (fotografia 23). Takie rozłożenie sprzyja równomiernemu rozkładowi światła i redukuje ryzyko powstawania lokalnych przekłamań barwnych. Mało tego: elektronika lampy – naszpikowana sensorami odległości (dalmierzami) i oświetlenia – dynamicznie reaguje na to, co dzieje się w polu operacyjnym i wokół niego, automatycznie dostosowując jasność czy też... ściemniając moduły, które w danym momencie są zasłonięte np. przez głowę chirurga lub asystenta.



Fotografia 24. Oświetlacz endoskopowy LV-400LED marki Vimex (<http://t.ly/wTR34>)

Warto dodać, że diody LED trafiły nawet do tych obszarów medycyny, które jeszcze do niedawna wydawały się być całkowicie zdominowane przez inne źródła światła. Mowa o oświetlaczach do endoskopów i laparoskopów – te niepozorne urządzenia odgrywają niezwykle ważną rolę na salach zabiegowych, odpowiadają bowiem za właściwą iluminację ciasnych przestrzeni wewnątrz ciała pacjenta podczas zabiegów chirurgii małoinwazyjnej, artroskopii, gastroskopii i innych procedur tego typu. Oświetlacz LV-400LED polskiej marki Vimex (fotografia 24) oferuje silne światło o strumieniu 2650 lumenów, porównywalne do światła 180-watowej lampy ksenonowej (!). Temperatura barwowa wynosi 5700 K, a współczynnik CRI jest równy 95. Sterownik umożliwia regulację jasności w zakresie od 1 do 100%, a całe urządzenie pobiera moc na poziomie zaledwie 115 VA.

Podsumowanie

W artykule zaprezentowaliśmy szerokie spektrum konsumenckich, komercyjnych oraz profesjonalnych aplikacji najnowocześniejszych diod LED. Nie ulega wątpliwości, że jest to zaledwie reprezentatywny wycinek spośród tysięcy zastosowań, z którymi spotykamy się co rusz w naszym codziennym życiu. W domu, szkole, pracy, sklepie, klubie muzycznym, na koncercie czy w szpitalu – półprzewodnikowe źródła światła dominują nad wszystkimi innymi, konkurencyjnymi technologiami i nic nie wskazuje na to, by cokolwiek miało zatrzymać tę LED-ową rewolucję.

inż. Przemysław Musz, EP



Moduły LED COB

– rodzaje i aspekty aplikacyjne

Moduły LED COB zrobiły zawrotną karierę na rynku oświetleniowym. To wydajne, kompaktowe i nowoczesne źródła światła, które w rzeczywistości stanowią połączenie znanych od lat technologii. Warto poświęcić im jednak nieco więcej uwagi.

COB to skrót od słów Chip-On-Board, czyli dosłownie „chip na płytce”. Któż z nas nie spotkał się w swojej praktyce z tym określeniem? Technologia COB jest wprawdzie używana w wielu obszarach mikroelektroniki, lecz zdecydowanie najczęściej skrót ten bywa stosowany w odniesieniu do modułów LED. Budowę typowego komponentu LED COB widać na **rysunku 1**: dziesiątki, a nawet setki niewielkich diod chipowych umieszczane są na wspólnym podłożu i łączone techniką znaną z układów scalonych, czyli tzw. bondingiem, cienkimi drucikami – zarówno pomiędzy sobą, jak i pomiędzy wyprowadzeniami skrajnych diod a polami kontaktowymi służącymi do podłączenia zasilania.

Zaraz zaraz... ale czym różni się moduł COB od zwykłej płytki obsadzonej dużą ilością dyskretnych diod LED SMD? Pierwsza i najważniejsza różnica to gęstość upakowania diod – jest ona bowiem dużo większa, niż byłoby to możliwe do osiągnięcia przy zastosowaniu zwykłych, obudowanych diod do montażu powierzchniowego. Duża gęstość ułożenia elementów wymusza z kolei doskonale chłodzenie, gdyż popularne moduły COB często pracują z mocą sięgającą kilkudziesięciu, a nawet stu watów. Dlatego też podstawowym elementem konstrukcyjnym jest tu już nie zwykły laminat FR4, ale dość gruba, aluminiowa blacha, na której umieszcza się warstwę izolacyjną, będącą właściwym podłożem chipów oraz głównych padów połączeniowych.

Co ciekawe, bardzo rzadko w praktyce spotykane są moduły jednobarwne (**fotografia 1**) lub emitujące światło w dwóch pasmach (**fotografia 2**) – prawie zawsze technologia COB jest stosowana do produkcji źródeł białych, co ma zresztą duży sens praktyczny. Pokrycie całości wspólnym luminoforem pozwala zwiększyć efektywność

optyczną, bowiem poszczególne chipy iluminują (światłem niebieskim) nie tylko małą warstwę luminoforu pokrywającą bezpośrednio sam chip, ale także dość spory obszar materiału znajdującego się pomiędzy sąsiadującymi diodami (**fotografia 3**).

Dość rzadko (choć nie są to całkiem odosobnione przypadki) zdarza się, by producent na tym samym podłożu, na którym znajduje się właściwa część emitera COB, umieszczał dodatkowe drivery czy układy zasilające – na rynku można jednak znaleźć i takie rozwiązania (**fotografia 4**). Ich niewątpliwą zaletą stanowi łatwość podłączenia, gdyż nie wymagają one żadnych dodatkowych układów zewnętrznych – wystarczy podłączyć przewód zasilania i et voilà – mamy światło!

W swoich założeniach technologia COB jest raczej dosyć prosta i bardzo elastyczna, jeśli idzie o możliwości tworzenia różnych kształtów oraz rozmiarów modułów LED. Dlatego nie są już dziś zaskoczeniem „dziwne” moduły, które można za śmiesznie małe pieniądze nabyć na chińskich platformach sprzedażowych – na przykład okrągłe z dwoma koncentrycznymi pierścieniami, co widać na **fotografii 5**. Nie zaskakują też ogromne rozmiary czy nietypowe proporcje modułów – to zjawisko dobrze pokazuje **fotografia 6**.



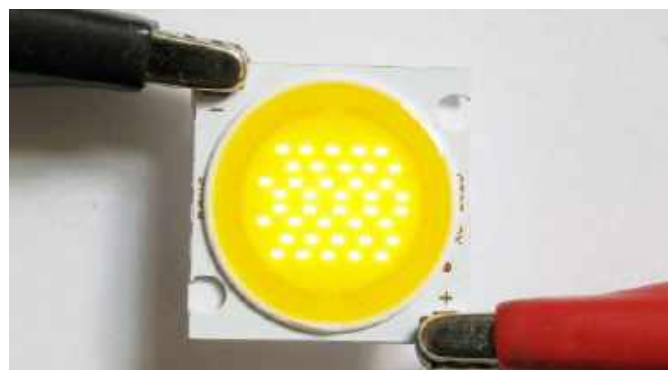
Fotografia 1. Zielony moduł COB (za: <http://t.ly/OFxD4>)



Rysunek 1. Budowa modułu LED COB (za: <http://t.ly/us0Jf>)



Fotografia 2. Specjalny moduł COB opracowany do wspomagania upraw (za: <http://t.ly/L27aE>)



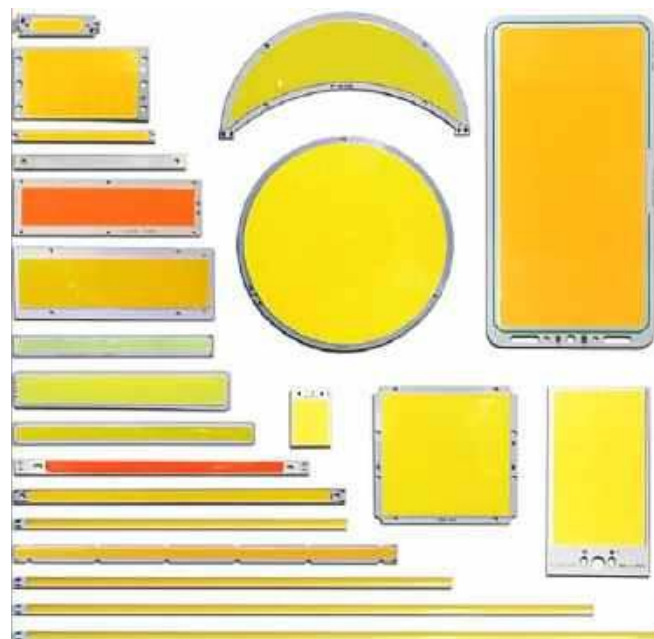
Fotografia 3. Moduł COB zasilany niskim napięciem – jak na dłoni widać poszczególne chipy LED (za: <http://t.ly/rX2ka>)



Fotografia 4. Ciekawe moduły COB – oprócz samego emitera na płytce jest też wbudowany zasilacz sieciowy (za: <http://t.ly/S0goh>)



Fotografia 5. Nietypowe, chińskie moduły COB (za: <http://t.ly/lqxy5>)



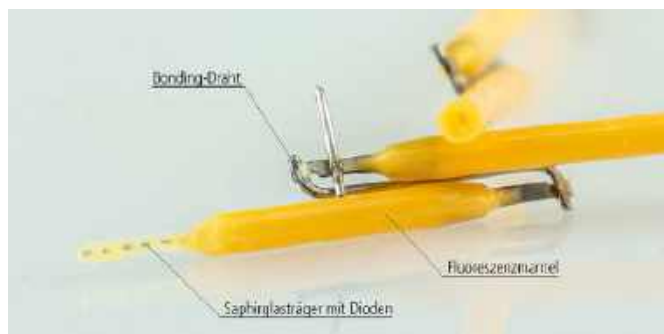
Fotografia 6. Bogactwo chińskich modułów COB nie zna granic (za: <http://t.ly/oiBeK>)

Filamenty LED

W pewnym momencie rozwoju nowoczesnych żarówek LED okazało się, że zwykle rozwiązania bazujące na diodach dyskretnych (np. takie, jak pokazano na **fotografii 7**) to troszkę za mało... Wprowadzie podobnie „sztańpowe” konstrukcje też świetnie sprawdzały się w praktyce, ale trzeba było opracować coś, co przyciągnie użytkowników do stosowania ledowych źródeł światła. Tak powstała żarówka filamentowa – żeby było śmieszniej, wzorowana na... pierwowzorze opracowanym w XIX wieku przed Thomasa Edisona (sic!). Ów „filament” to tak naprawdę nic innego, jak pewna odmiana COB, określana często mianem COG (Chip-on-Glass). Zastosowanie szkła jako materiału bazowego pozwala lepiej rozpraszać światło we wszystkich kierunkach. Reszta jest już niemal taka sama, jak w przypadku zwykłych modułów COB: szereg struktur diodowych zamontowany zostaje na wspólnym podłożu, a dwa wyprowadzenia (tym razem w postaci cienkich blaszek) wystają poza wspólną „zalewę” wykonaną z luminoforu (co widać jak na dłoni na **fotografii 8**). Tak zbudowany filament ma pewną wadę – jest sztywny i kruchy, zupełnie nie przypomina pod tym względem klasycznego żarnika wolframowego. W przypadku niektórych żarówek nie stanowi to jednak żadnego problemu, bo poszczególne „włókna” i tak są montowane



Fotografia 7. Typowa żarówka LED złożona z dyskretnych diod (za: <http://t.ly/6N8ph>)



Fotografia 8. Anatomia filamentu LED (za: <http://t.ly/0Zr1L>)



Fotografia 9. Żarówki LED w stylu vintage (za: <http://t.ly/hT1jn>)

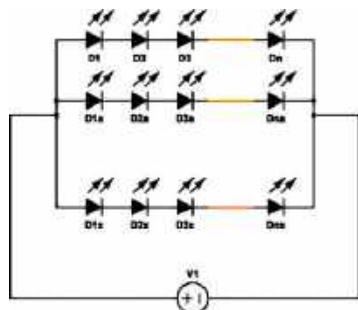


Fotografia 10. Zakończenia elastycznych filamentów (za: <http://t.ly/hNWk4>)



Fotografia 11. Elastyczny filament o barwie niebieskiej (za: <http://t.ly/Oig9n>)

na stałe w szklanej bańce, gdzie nie grożą im żadne większe naprężenia. Efekt jest naprawdę ciekawy i chętnie używany do tworzenia żarówek stylizowanych na lata rewolucji elektrycznej XIX wieku (**fotografia 9**). Takie konstrukcje, natychmiast kojarzące się z pierwszymi, nieporadnymi żarówkami sprzed wielu dekad, nazywa się (a jakże!) mianem „żarówek Edisona”. Słynny wynalazca byłby pewnie z tego dumny.



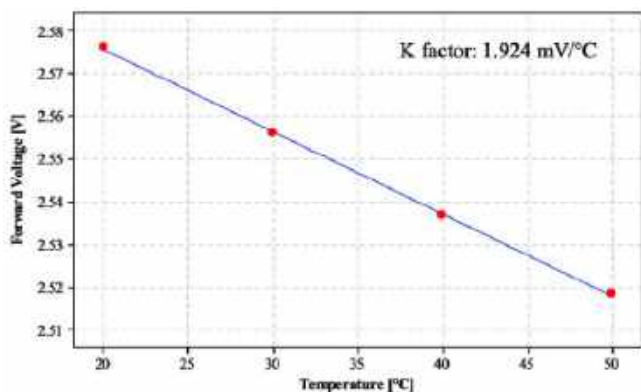
Rysunek 2. Typowe połączenia chipów LED w module COB (za: <http://t.ly/Q62yG>)

Na szczęście z biegiem lat powstały też filamenty elastyczne, w których niewygodny element (kruche szkło) został zastąpiony giętkim podłożem. Tego typu elementy można już bez problemu kupić i to za niewielkie pieniądze (**fotografia 10, fotografia 11**).

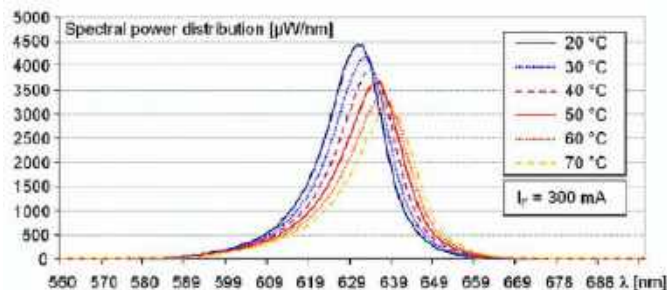
A co z zasilaniem nowych „tworów”? W tym miejscu należy dodać, że filamenty dostępne na rynku występują w wielu różnych odmianach. W niektórych diody połączone zostają szeregowo. Ma to sens zwłaszcza w przypadku filamentów stosowanych w żarówkach sieciowych: im więcej diod, tym wyższe napięcie przewodzenia (sumaryczne), a dzięki temu układ zasilania ma „mniej roboty” z obniżaniem napięcia 230 V (lub 115 V) do odpowiedniego poziomu. Część filamentów – zwłaszcza te sprzedawane detalicznie amatorom DIY – bywa wykonywana w wersji równoległej, przez co napięcie przewodzenia oscyluje wokół 3 V.

W przypadku modułów COB sprawa okazuje się bardziej złożona. Liczba diod znajdujących się w pojedynczym takim elemencie bywa ogromna, przez co ani połączenie szeregowe, ani równoległe, nie miałyby praktycznie żadnego sensu. Dlatego powszechnie stosuje się konfiguracje mieszane (patrz **rysunek 2**), w efekcie których sumaryczne napięcie oraz całkowity prąd zasilania znajdują się w racjonalnych granicach. Typowe napięcia zasilania modułów LED COB wynoszą kilkadziesiąt woltów, zaś prąd – od kilkuset miliamperów do kilku amperów.

Uwaga! Bardzo ważne jest prawidłowe zasilanie modułu COB. W przypadku „gołych” modułów (bez wbudowanej elektroniki) zalecane jest stosowanie zasilacza stałoprądowego. A dlaczego? Sprawa jest bardzo prosta, a wyjaśnienia należy szukać w zachowaniu półprzewodnika. Otóż napięcie przewodzenia diod LED (V_f) spada wraz ze wzrostem temperatury struktury. Jeżeli do COB podłączymy źródło napięcia z niewielką rezystancją szeregową, to przepływ prądu spowoduje delikatne podgrzanie struktury. To zaś doprowadzi do obniżenia V_f , co zwiększy spadek napięcia na rezystorze, a w efekcie podwyższy natężenie przepływającego prądu... i tak w kółko, aż do przegrzania modułu. Efekt ten jest szczególnie



Rysunek 3. Zależność napięcia przewodzenia diody LED od temperatury (za: <http://t.ly/NO7HM>)



Rysunek 4. Zależność spektralnego rozkładu mocy od temperatury diody LED (za: <http://t.ly/mBcHi>)



Fotografia 12. Wysokiej klasy moduły LED COB o wysokiej wartości CRI (Color Rendering Index – współczynnik oddawania barw, za: <http://t.ly/uxsX2>)

niebezpieczny w przypadku COB, gdyż nawet zmiana V_f o pojedyncze miliwoltę skumuluje się w wyniku szeregowego połączenia diod w poszczególnych gałęziach. Dlatego jedyną bezpieczną metodą zasilania COB jest właśnie użycie źródła prądowego, które nie dopuści do termicznej destabilizacji naszego modułu.

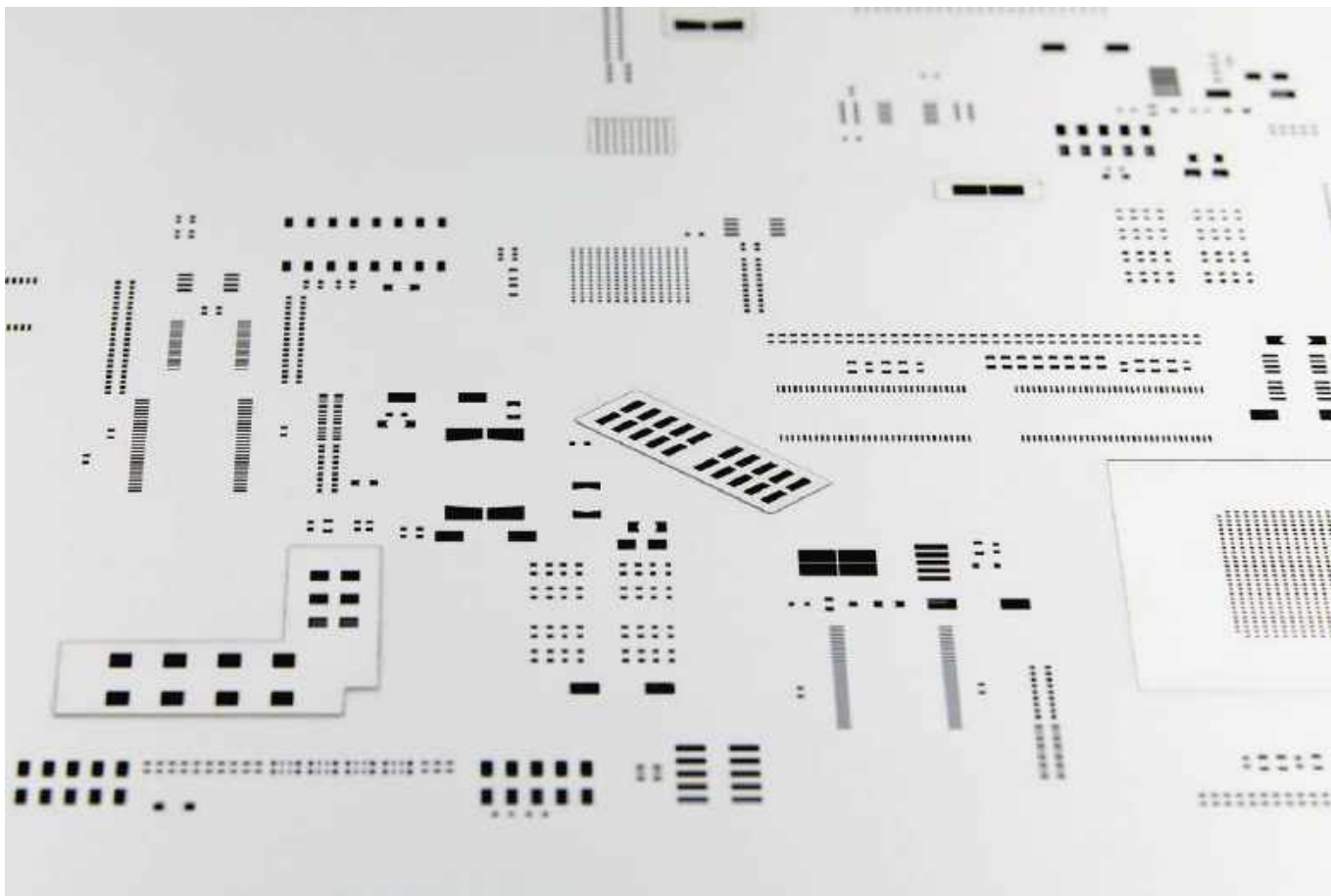
Kwestie termiczne

Ochrona diod LED (w tym także COB) przed nadmierną temperaturą jest konieczna przede wszystkim ze względu na ryzyko przegrzania, prowadzącego do skrócenia ich żywotności albo nawet gwałtownego przepalenia. Ale to nie wszystko. Temperatura wpływa także na widmo światła emitowanego przez diody – na **rysunku 4** widać, doskonale że pik emisji przesuwają się w miarę wzrostu temperatury. Niby niewiele (około 10 nm przy podgrzaniu o 50°C), ale w niektórych aplikacjach może to mieć znaczenie – zwłaszcza jeżeli w grę wchodzi jeszcze inne czynniki, np. wrażliwość luminoforu diod białych. Badania dowodzą, że wraz ze wzrostem temperatury parametry diody zaczynają się „rozjeżdżać”: zmieniają się amplitudy, szerokości oraz położenie pików widma. Jeżeli problem dotyczy wysokiej klasy diod określanych jako full spectrum (czyli emitujących światło o składzie spektralnym naśladującym światło słoneczne – przykład widać na **fotografii 12**), to zbyt wysoka temperatura zniweczy wysiłki projektantów, mające na celu dokładne odwzorowanie widma idealnie białego. Dlatego też projektując urządzenie z modułem COB należy zawsze skrupulatnie przemyśleć (i przetestować!) odpowiedni system chłodzenia – efektywny przykład widać na **fotografii 13**.

Jakub Nowicki



Fotografia 13. Moduł LED COB z ogromnym radiatorem i zintegrowanym układem optycznym (za: <http://t.ly/cjkBl>)



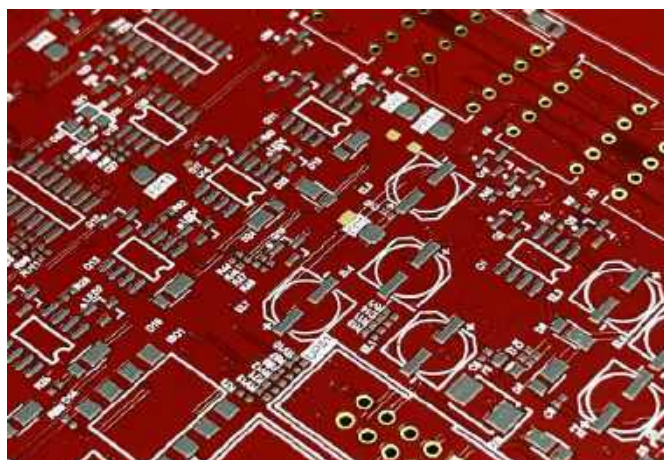
Szablony SMT, co warto wiedzieć

Jakość wykonania szablonu do nakładania pasty lutowniczej ma zaskakująco silny wpływ na efekty montażu. Okazuje się, że defekty związane z samym tylko nanoszeniem pasty lutowniczej mogą odpowiadać za nawet 60% błędów montażowych w technologii SMT. Dlatego właśnie dbałość zarówno o stabilność procesu sitodruku pasty lutowniczej, jak i jakość oraz właściwy projekt szablonu SMT stanowi czynnik, którego znaczenie trudno przecenić – zwłaszcza w produkcji wysokonakładowej.

Produkcja szablonów SMT

Grubość szablonu oraz rozmiary wycinanych laserowo apertur określają ilość pasty lutowniczej nałożonej na pady PCB. Kontrolę ilości pasty znajdującej się na powierzchni poszczególnych pól lutowniczych realizuje się z użyciem specjalistycznych maszyn pomiarowych (SPI) lub też odpowiednio wyposażonych stanowisk automatycznej inspekcji optycznej (AOI).

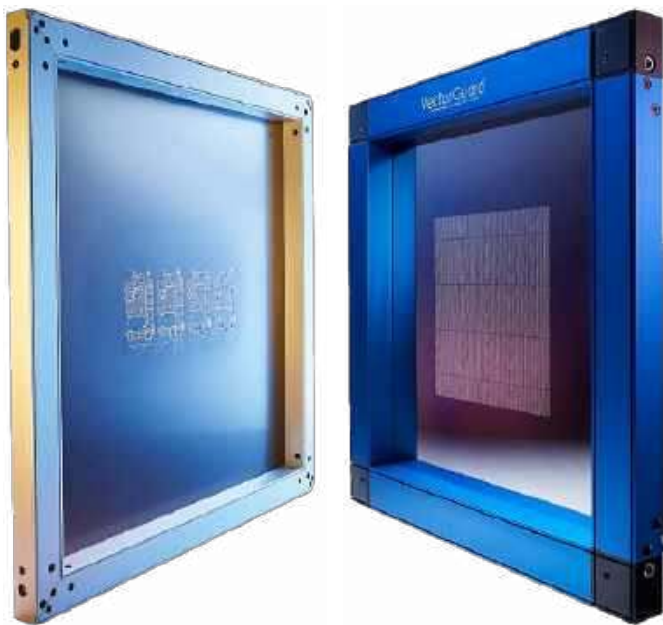
Szablony SMT są zwykle produkowane z arkuszy blachy ze stali nierdzewnej o grubości w zakresie od 80 do 180 μm (typu SS301), o strukturze drobnoziarnistej 2...3 μm , określanej mianem fine grain. Odpowiednio mały rozmiar ziaren stopu pozwala bowiem ograniczyć erozję laserową brzegów wycinanych apertur, sprzyjając tym samym łatwemu i jednolitemu uwalnianiu pasty lutowniczej. Jeżeli wymogi miniaturyzacji docelowej PCB narzucają szczególne ograniczenia w zakresie rastra, w praktyce stosuje się folie z elektrofornowanego niklu o grubości 100...150 μm , o jeszcze drobniejszej



ziarnistości, a co za tym idzie – lepszym uwalnianiu pasty. Zakres możliwości cięcia lasera pozwala na wykonywanie szablonów i detali z folii stalowych SS304 o grubości 50...1000 μm .

Park maszynowy firmy Semicon dysponuje trzema najwyższej klasy urządzeniami laserowymi firmy LPKF, w tym dwoma G6080 (wyposażonymi w moduł mikrospawania w osłonie azotu) oraz jednym G6060.

Układy pozycjonowania głowic laserowych ww. obrabiarek zapewniają dokładność do 2 μm , co potwierdzone jest okresową kalibracją urządzeń, wykonywaną przez serwis producenta. Warto dodać, że obszar roboczy wycinanych szablonów wynosi standardowo 600×800 mm, ale w razie potrzeby (np. do celów produkcji taśm LED) możliwe jest wykonanie szablonów o długości nawet do 1800 mm.



Dostępne grubości szablonów, w zależności od zastosowanego materiału, wynoszą:

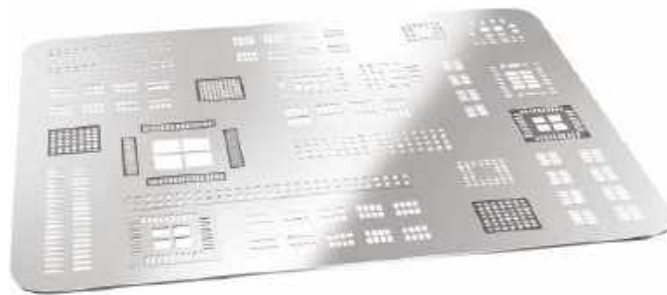
- folia stalowa SS304: 50...1000 μm ,
- folia fine grain SS301: 80...180 μm ,
- elektroformowany nikiel: 100...150 μm .

Autorska metoda optymalizacji drogi cięcia, opracowana w firmie Semicon, zapewnia minimalizację przeciążenia termicznego materiału, zapobiegając lokalnym odkształceniom – szczególnie w obszarach szablonu o dużym zagęszczeniu układów drobnorastrowych. System kamer pozycjonujących głowicę lasera umożliwia natomiast optymalne skupienie wiązki, zmierzenie geometrii wyciętych otworów oraz uzupełnienie wykonanego wcześniej szablonu.

Standardowy czas realizacji zamówienia to 2 dni

Semicon podejmuje się także produkcji i dostawy szablonów w trybie ekspresowym – nawet w dniu dostarczenia plików. Preferowane pliki to gerbery w formatach RS274X i RS274D. W razie potrzeby istnieje możliwość odpowiedniej modyfikacji dostarczonej przez klienta dokumentacji produkcyjnej.

Należy pamiętać, że finalna jakość połączeń lutowanych zależy nie tylko od wyciętych w szablonie otworów, ale także od ilości



nałożonej za ich pomocą pasty lutowniczej. W przypadku płytek drukowanych, na których montowane są komponenty o dużym zróżnicowaniu rozmiaru apertur, zalecane jest stosowanie szablonów stopniowanych (tj. o zróżnicowanej grubości folii stalowej), które Semicon wytwarza za pomocą dwóch metod.

1. Szablony mikrospawane

Mikrospawanie polega na precyzyjnym wycięciu otworów w bazowym szablonie i następnie wspawaniu wyciętych uprzednio formatków z blachy (wykonanej z tego samego materiału) o odpowiedniej grubości. Maksymalna różnica grubości w tej metodzie to 80 μm , a minimalna odległość pozycjonowania elementów SMD od uskoku grubości wynosi 200 μm .

2. Szablony frezowane

Frezowanie wyciętych laserowo szablonów pozwala uzyskać dokładność grubości szablonu do 5 μm , a różnicę poziomów do nawet 400 μm . Tak wykonane szablony stopniowane stanowią technologiczną nowość i pozwalają na niezawodny montaż – na jednej płycie – zarówno najmniejszych komponentów dyskretnych czy układów scalonych o drobnym rastrze, jak i elementów o dużych rozmiarach (np. dławików dużej mocy). Firma Semicon dysponuje wieloletnim doświadczeniem w zakresie projektowania i modyfikacji tego typu szablonów.

Szablony SMT z nanopowłokami

Dalsza optymalizacja procesu uwalniania pasty z apertur jest możliwa przy zastosowaniu nanopowłok nakładanych na powierzchnię szablonu. Dostępne są przy tym dwie techniki:

1. Nanoszenie warstwy nanomateriału za pomocą nasączonej aktywną powłoką ściereczki.
2. Permanentne nanoszenie powłok metodami CVD.

REKLAMA

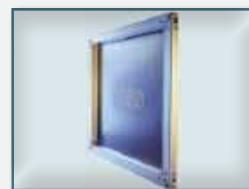


Innowacyjne produkty
Innowacyjne technologie

Szablony SMT

- Wycinanie szablonów
- Szablony mikrospawane
- Szablony frezowane
- Nanopowłoki

- Szablony do układów BGA
- Naciągi:
 - pneumatyczne - zasilane ramy samonapinające
 - aluminiowe z naciągiem siatkowym
 - ramy w standardzie Vector Guard





Sama powłoka to w istocie cienka warstwa fluoro-polimeru o grubości 1...4 μm , pozwalająca na obniżenie tzw. współczynnika Area Ratio do 0,5, co skutkuje lepszym transferem pasty na pad lutowniczy, w porównaniu ze zwykłymi szablunami pozbawionymi powłoki. Nanomateriał umożliwia także poprawę parametrów drukowania, szczególnie w przypadku płyt o gęstym upakowaniu drobnych komponentów.

Stabilność naciągu szablonu gwarancją jakości procesu montażowego

Wielkość oraz stabilność naciągu szablonu także wpływa na jakość procesu lutowania SMT. Firma Semicon oferuje trzy metody w tym zakresie.

1. **Pneumatyczne, zasilane ramy samonapinające** w standardach Zelfelx oraz Quatroflex – wymagają okresowej kontroli siły naciągu, przez co są stosowane zwykle do małych i średnich serii produkcyjnych.
2. **Szablony na ramach aluminiowych z naciągiem siatkowym** sprawdzają się w przypadku dużych i wielkich partii PCB,

pozwalają bowiem na wykonanie nawet do 200 tysięcy cykli zadruku pasty przy zachowaniu powtarzalnych parametrów nanoszenia lutowni na pady. Ograniczeniem jest w tym przypadku grubość używanych ram aluminiowych (zwykle 38 mm), co wymaga znacznych powierzchni odstawczych.

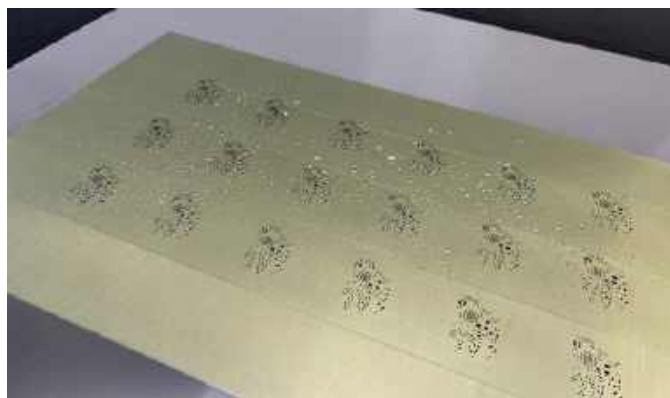
3. **Ramy w standardzie Vector Guard** – Semicon jest oficjalnym franczyzobiorcą licencji Vector Guard i oficjalnym partnerem licencjodawcy firmy ASM (dawny DEK). Ramy Vector Guard zapewniają stabilność naciągu poprzez zastosowanie mechanicznego systemu i sprawdzają się także w przypadku niewielkich aplikacji z małymi rastrami elementów (opcja high tension). Znormalizowana grubość szablonu Vector Guard to 8 mm, co pozwala ograniczyć zapotrzebowanie na powierzchnię odstawczą dla szablony SMT. Warto dodać, że popularność tego typu ram sięga już teraz do 75% aplikacji na zachodzie Europy i stale rośnie.

Sita do układów w obudowach BGA

Firma Semicon projektuje i wytwarza szablony (sita) do napraw układów BGA. Umożliwiają one reballing, czyli rekonstrukcję kulek za pomocą pasty. W ofercie firmy znajdują się ponadto usługi modyfikacji szablony z „pasty na klej” oraz grawerowania i wykonywania znaczników.

W przypadku braku oryginalnej dokumentacji technicznej możliwe jest również przeprowadzenie procesu inżynierii odwrotnej mikroszablonu do nakładania kulek BGA lub pasty – na podstawie fizycznego egzemplarza układu scalonego, dostarczonego przez klienta.

Firma Semicon dysponuje zasobem i stałym dostępem do folii z różnych stopów stali, takich jak: inconel, kovar, nikiel, tantal czy molibden, świadcząc zarazem usługi wycinania detali mechanicznych. W ofercie znajdują się również szafy do przechowywania szablony, rolki czyszczące, płyny myjące oraz opakowania ESD.



Testery funkcjonalne FCT – niezawodne rozwiązania w testowaniu elektroniki

W dzisiejszym świecie elektroniki niezawodność i jakość produktów mają fundamentalne znaczenie. Do zapewnienia wysokiej jakości urządzeń niezbędne jest przeprowadzenie dokładnych testów funkcjonalnych na każdym etapie produkcji. W tym celu stosuje się testery funkcjonalne (FCT) – ich rola jest nieoceniona, gdy chcemy upewnić się, że końcowy produkt spełnia wszystkie wymagania techniczne oraz normy jakościowe.

Czym są testery funkcjonalne (FCT)?

Testery funkcjonalne (Functional Circuit Testers, FCT) to zaawansowane urządzenia służące do weryfikacji poprawności działania gotowego produktu elektronicznego lub poszczególnych jego modułów. FCT przeprowadzają testy na poziomie funkcjonalnym, symulując rzeczywiste warunki pracy urządzenia. Dzięki temu możliwa jest całkowicie automatyczna praca testera, począwszy od włączenia testowanej elektroniki i zmierzenia parametrów jej zasilania, poprzez programowanie JTAG lub podobne, aż po testowanie funkcji oraz kalibrację urządzenia.

W zależności od potrzeb, testery mogą być wyposażone w oprogramowanie skryptowe lub oparte na środowisku LabVIEW, co pozwala na elastyczne dostosowanie procesu testowania do specyficznych wymagań produkcyjnych. W niektórych przypadkach mobilne testery autonomiczne, które pracują bez udziału zewnętrznego komputera PC, są używane do zwiększenia elastyczności linii produkcyjnej.

Zastosowanie testerów funkcjonalnych z fiksturami

Testery funkcjonalne z fiksturami (fixture-based testers) są szczególnie efektywne w środowiskach produkcyjnych, gdy liczy się czas i precyzja. Fikstury to specjalnie zaprojektowane platformy, na których umieszcza się testowane urządzenia. Pozwalają one na automatyczne podłączenie wszystkich niezbędnych sygnałów, zasilania oraz interfejsów pomiarowych do badanego produktu – taki zabieg minimalizuje ryzyko błędów ludzkich, a także skracza czas potrzebny na przeprowadzenie testów.

Testery z fiksturami są często wyposażone w zaawansowane systemy weryfikacji, które mogą realizować takie procesy, jak testowanie wysokoprądowe, pomiary sygnałów analogowych oraz cyfrowych, czy ich zadawanie. Wbudowana szafa sterownicza ma wbudowane specjalizowane, automatyczne przełączniki sygnałowe

zbudowane wg autorskiego projektu, może być również wykonana z komponentów handlowych. Możliwość dodatkowej integracji z oprogramowaniem ułatwia pełną automatyzację procesu testowania, w tym analizę wyników i generowanie raportów.

Korzyści płynące z zastosowania testerów funkcjonalnych

Zastosowanie testerów funkcjonalnych, zwłaszcza tych z fiksturami, niesie ze sobą wiele korzyści. Automatyzacja procesu testowania pozwala na znaczne zwiększenie efektywności kontroli jakościowej, co okazuje się szczególnie ważne w przypadku produkcji masowej. Automatyczne podłączanie sygnałów i zasilania eliminuje ryzyko błędów typowych dla ręcznej obsługi, a to z kolei przekłada się na większą precyzję oraz powtarzalność wyników. Dzięki skróceniu czasu testowania możliwa jest redukcja kosztów operacyjnych, a to poprawia ogólną wydajność produkcji. Co więcej: testery funkcjonalne są skalowalne i łatwo dostosować je do różnych produktów, dlatego stanowią one uniwersalne narzędzie w każdej fabryce elektroniki.

Programowanie układów scalonych przed montażem SMT

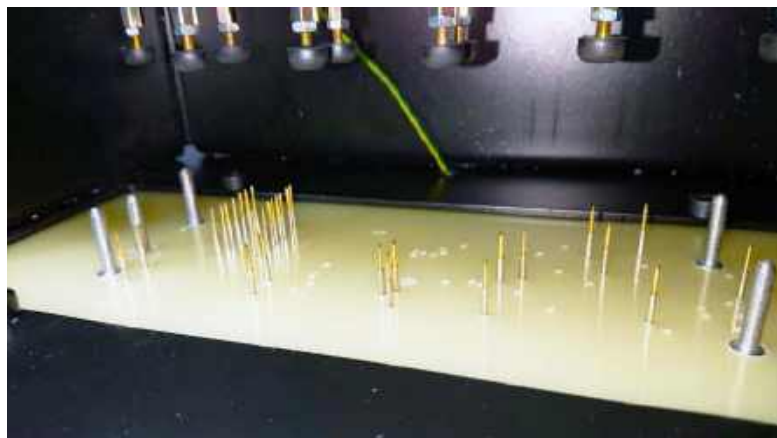
Sporadycznie układy scalone muszą być zaprogramowane przed wlutowaniem w płytkę. Czasami wynika to z nietypowego algorytmu programowania, uproszczeń lub zabezpieczeń na gotowej PCB czy braku odpowiedniego bootloadera w ROM. Programowanie odbywa się zazwyczaj w specjalnie do tego przygotowanych podstawkach, podłączanych pod odpowiednie programatory. Jeśli wgrywanie firmware'u do pamięci wymaga specjalnych algorytmów, które nie są obsługiwane przez typowe programatory, to należy je przygotować indywidualnie pod konkretny układ.

Podsumowanie

Testery funkcjonalne z fiksturami są niezbędnym narzędziem w nowoczesnej produkcji elektroniki. Można za ich pomocą nie tylko zwiększyć efektywność procesu testowania, ale także zapewnić, by każdy produkt opuszczający linię produkcyjną spełniał najwyższe standardy jakości. Dzięki zaawansowanym funkcjom automatyzacji, a także możliwości integracji z różnymi systemami oprogramowania testery FCT oferują kompleksowe rozwiązania dla każdej linii produkcyjnej, a tym samym przyczyniają się do lepszej jakości produktów i większej wydajności produkcji. Możliwość wykonania całości testera w jednej firmie pozwala zaoszczędzić czas na jego realizacji, a programowanie układów scalonych przed produkcją sprzyja optymalizacji procesu logistycznego.

REKLAMA

Testery FCT • Programowanie układów • Kompleksowo



AWICAM
EMBEDDED HARDWARE
& SOFTWARE HOUSE

Produkcja płytek drukowanych – przewodnik w pigułce

Płytki drukowane to kręgosłup (jeśli nie cały szkielet) współczesnej elektroniki. Stanowią one podstawowy element niemal każdego urządzenia elektronicznego, od smartfonów czy zaawansowanych urządzeń medycznych po sterowniki pralek czy melodyjkę w kartce pocztowej. PCB zapewniają połączenia (elektryczne, ale i mechaniczne) dla setek czy nawet tysięcy elementów elektronicznych. Pośredniczą między skomplikowanymi układami elektronicznymi, umożliwiając efektywną komunikację i zasilanie wszystkich podzespołów. Bez PCB dzisiejszy świat technologii oraz automatyzacji nie mógłby istnieć w takiej formie, jaką znamy. W tym artykule przyjrzymy się, w jaki sposób są one projektowane, a także produkowane.

Nie ulega wątpliwości, że nie można wyobrazić sobie współczesnej elektroniki bez płytek drukowanych. Nie jest to oczywiście jedyna technologia łączenia elementów elektronicznych, ale z pewnością pozostaje najpopularniejsza. Historia PCB sięga pierwszej połowy XX wieku, a dokładniej lat 30., gdy austriacki inżynier, Paul Eisler, użył cienkiej warstwy metalu do zastąpienia najpopularniejszych wtedy połączeń drutowych. Rozwiązanie to było stosowane głównie w urządzeniach radiowych. Jak łatwo zrozumieć, II wojna światowa znacznie przyspieszyła rozwój tej technologii, a co za tym idzie – przysporzyła popularności PCB.

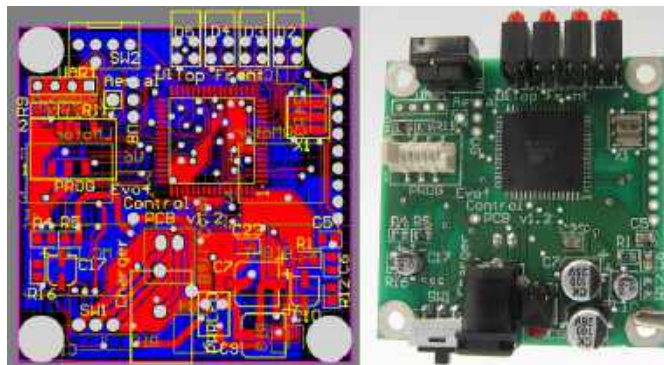
Dzięki postępowi w technologii materiałowej PCB zaczęły pojawiać się po wojnie również w zastosowaniach komercyjnych. Rozwój tranzystorów, a następnie układów scalonych w kolejnych dekadach jeszcze wyraźniej przyspieszył popularzację PCB, które stawały się coraz bardziej złożone, wielowarstwowe, ale także kompaktowe. Pozwoliły one na miniaturzację urządzeń. Powstanie kolejnych wynalazków – takich jak soldermaska (maska przeciwlutowa) – i zastosowanie automatycznego montażu elementów, w połączeniu z technologią PCB umożliwiło również redukcję kosztów i uruchomienie produkcji przemysłowej urządzeń elektronicznych.

W artykule prześledzimy cały proces, pozwalający zaistnieć współczesnej płytce drukowanej – od projektowania, poprzez przygotowanie dokumentacji, aż po finalną produkcję.

Proces produkcji PCB Projektowanie i przygotowanie plików (Design & Gerber Files)

Pierwszym krokiem do stworzenia płytki drukowanej jest oczywiście jej zaprojektowanie. Proces ten zaczyna się na ogół od stworzenia schematu ideowego. Inżynier, korzystając z oprogramowania EDA (ang. Electronic Design Automation), umieszcza na schemacie symbole reprezentujące komponenty elektroniczne i łączy je ścieżkami. Jednocześnie każdy z elementów na schemacie jest powiązany ze zdefiniowanym modelem obudowy, umieszczanym w projekcie PCB. Edytor schematów generuje listę połączeń (tzw. netlistę), która stanowi podstawę do projektowania PCB.

Początkowo płytki projektowało się analogowo – bez komputera. Jednak nawet gdy zdigitalizowano już rysowanie schematów, to skomplikowanie mozaiki PCB nie pozwalało na odejście od sprawdzonych metod – tworzenia projektu przy użyciu kalki



Fotografia 1. PCB w programie EDA (po lewej) i gotowa płytka drukowana (po prawej)

lub folii mylarowej i specjalnych naklejanych taśm oraz szablonów. Taśmy te dostępne były w kształtach ścieżek o różnej grubości oraz w kształcie pól lutowniczych (należy pamiętać, że mowa o czasach przed wprowadzeniem technologii montażu SMD, jak i przed popularzacją układów scalonych, więc układy pól były na ogół proste) – w kolorach (najczęściej) niebieskim i czerwonym (i dlatego we współczesnym oprogramowaniu EDA tych kolorów używa się do oznaczania ścieżek na warstwie dolnej i górnej płytki). Jeśli projekt PCB wymagał kilku warstw, składało się z kilku warstw folii, po jednej dla każdej warstwy na PCB. Folie te na ogół projektowane były w skali, a następnie pomniejszane optycznie w celu naświetlania klisz do produkcji PCB.

Dopiero w latach 80. XX wieku, wraz z rozwojem komputerów, zaczęły powstawać pierwsze pakiety EDA. Jednym z pionierów był program Protel (później przekształcony w znany Altium Designer), który znacznie uprościł i zautomatyzował proces projektowania. Wprowadzenie nowoczesnych pakietów EDA pozwoliło na tworzenie coraz bardziej złożonych projektów, otwierając tym samym drogę do rewolucji w elektronice.

Obecnie na rynku dostępnych jest wiele programów do projektowania płytek drukowanych – Altium Designer, KiCad, Eagle, EasyEDA, OrCAD, PADS, CircuitMaker czy LibrePCB. Pakiety te różnią się możliwościami i ceną (niektóre są darmowe, niektóre bardzo kosztowne), ale niniejszy artykuł nie jest miejscem do ich omawiania. Wspólna dla większości systemów EDA jest możliwość zdefiniowania zasad projektowania i weryfikacji, czy nasz projekt je spełnia (DRC – Design Rule Check). Jest to istotne, gdyż upewnia nas na etapie cyfrowego projektu, że płytka jest produkowalna w założony przez nas sposób. Inny wspólny element to możliwość eksportu projektu w postaci plików produkcyjnych.

Podstawowym formatem plików produkcyjnych jest tzw. Gerber, zawierający informacje niezbędne do opisu PCB. Pliki Gerber są używane od lat 60. XX wieku. Ich nazwa pochodzi od nazwiska ich wynalazcy – H. Josepha Gerbera. Obecnie są one powszechnie akceptowanym standardem w całym przemyśle i praktycznie wszystkie programy EDA potrafią generować pliki wyjściowe w tym formacie, a właściwie całą rodzinie formatów. Obecnie używane są trzy jego rodzaje:

- Standardowy Gerber (znany jako RS-274-D) – najstarsza forma plików Gerber, obecnie generalnie niewspierana.
- Rozszerzony Gerber (znany jako RS-274X), który jest najbardziej popularny.
- Gerber X2 – najnowszy format plików Gerber.

Niezależnie od wersji formatu, plik ten zawiera wektorową reprezentację obrazu na PCB. Zasadniczo jest to zbiór dwuwymiarowych współrzędnych, który opisuje każdą warstwę płytki. Oprócz ścieżek czy warstw napisów, pliki Gerber zawierają również informacje o warstwach soldermaski, a także pasty lutowniczej czy kleju. Jeśli korzystamy z innych, dodatkowych operacji, informacje o nich także można opisać za pomocą pliku Gerber. Producenci płytek opierają się na tych plikach, aby sterować urządzeniami produkcyjnymi, dlatego nawet drobne błędy mogą prowadzić do poważnych wad w gotowym produkcie. Projektanci muszą zatem zadbać o to, aby generowane pliki Gerber precyzyjnie odwzorowywały zamierzony projekt i były zgodne ze specyfikacją producenta.

Do najczęściej spotykanych błędów w pracy z plikami Gerber należą m.in.:

1. Brakujące lub nieprawidłowe warstwy

Brakujące lub niepoprawnie zdefiniowane warstwy stanowią jeden z najczęstszych i najpoważniejszych błędów w plikach Gerber. Może on wynikać z błędów podczas generowania plików lub niekompletnego eksportu danych z pakietu EDA. Brak tych warstw potencjalnie prowadzi do wyprodukowania wadliwych płytek.

W większości pakietów EDA wybieramy, z których warstw generowane są pliki Gerber. Należy dołożyć wszelkich starań, aby odpowiednie warstwy znalazły się w odpowiednim pliku: w osobnych plikach umieszczamy ścieżki każdej warstwy, soldermaskę, opisy na górnej i dolnej stronie PCB itd. Dodatkowe warstwy mogą obejmować m.in. złoczenie selektywne czy miejsce nakładania innych pokryć (np. konformalnego – specjalnego lakieru, który zabezpiecza płytkę i elementy na niej przed wpływem czynników atmosferycznych) na PCB.

2. Niewłaściwie wyrównane lub rozmieszczone warstwy

Nawet jeśli mamy już wszystkie warstwy, to mogą być one nieprawidłowo wyrównane względem siebie. Częstym błędem jest ich niewłaściwe rozmieszczenie, co wiąże się np. z błędami w trakcie projektowania, ale też ze złym wyrównaniem punktu zerowego dla różnych warstw, co z kolei może wynikać z nieprawidłowości przy generowaniu bądź konwersji plików. W efekcie może dojść do zwarcia czy wręcz przeciwnie – przerw w ścieżkach. Taka PCB z pewnością nie będzie nadawała się do użytku.

3. Niepoprawna definicja apertur

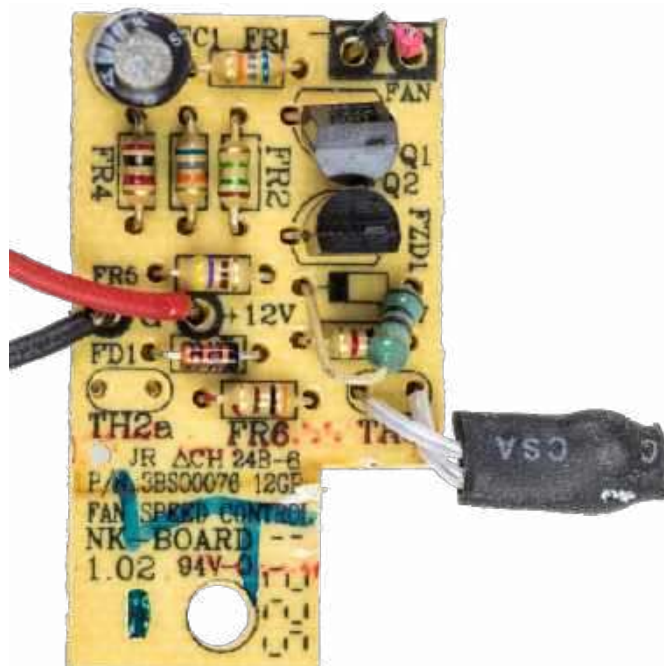
Apertury określają rozmiary i kształty elementów na PCB, takich jak pola lutownicze, przelotki czy ścieżki. Nieprawidłowe definicje apertur potencjalnie mogą skutkować niewłaściwymi rozmiarami lub kształtami tych elementów – a w konsekwencji przysporzyć problemów z montażem komponentów, lutowaniem lub integralnością sygnału.

4. Niezgodność jednostek miary

Pliki Gerber mogą używać różnych jednostek miary, na ogół milimetrów lub milsów. Niezgodność jednostek w zestawie plików może powodować problemy z wymiarami, co prowadzi do niewłaściwego przeskalowania PCB. Należy dopilnować, aby skala i jednostki wpisane w naszym pakiecie EDA przy generowaniu plików Gerber były takie, jakich spodziewa się producent PCB.

5. Nakładające się lub zduplikowane obiekty

Ostatnia uwaga związana jest częściowo z pierwszą. Jeśli w wybranych warstwach w oprogramowaniu EDA znajdować się będą niepoprawne obiekty – lub wybrane zostaną omyłkowo np. dodatkowe warstwy – spowodować może to nałożenie dodatkowych obiektów na PCB, co przyczyni się do powstania np. zwarcia.



Fotografia 2. Płytki PCB wykonane z laminatu FR1

Oprócz plików Gerber do produkcji płytki drukowanej potrzebne są jeszcze pliki definiujące otwory w PCB oraz, opcjonalnie, miejsca frezowania w płytce.

Najczęściej używanym formatem plików owierty PCB jest Excellon, nazwany tak na cześć maszyny, która była jedną z pierwszych automatycznych wiertarek stosowanych w produkcji płytek drukowanych. Excellon to standardowy format używany przez większość producentów – zasadniczo jest on plikiem tekstowym zawierającym zestaw współrzędnych (x, y), które wskazują dokładne położenie otworów, a także specyfikacje narzędzi (rozmiary wiertel). Format Excellon obsługuje ponadto dodatkowe parametry, m.in. głębokość wiercenia (choć w większości PCB otwory wierci się na standardową głębokość – przelotowo). Pliki te mogą również definiować różne typy owierty, np. zwykłe przelotki (vias) czy przelotki ślepe (blind vias) lub zagrzebane (buried vias), które łączą tylko wybrane warstwy płytki.

Pliki owierty korzystają z tego samego układu współrzędnych i systemu jednostek miary (milimetry lub cale), co pliki Gerber, aby uniknąć problemów z niezgodnością rozmiarów. Niewłaściwe skalowanie lub format pliku owierty może prowadzić do błędów (np. nieprawidłowego umiejscowienia otworów), co może skutkować problemami w montażu komponentów czy w ogóle powstaniem płytki niezdatnej do użycia.

REKLAMA



- obwody drukowane
jedno, dwustronne
oraz wielowarstwowe

- obwody drukowane
do zastosowań LED

- prototypy i serie
produkcyjne

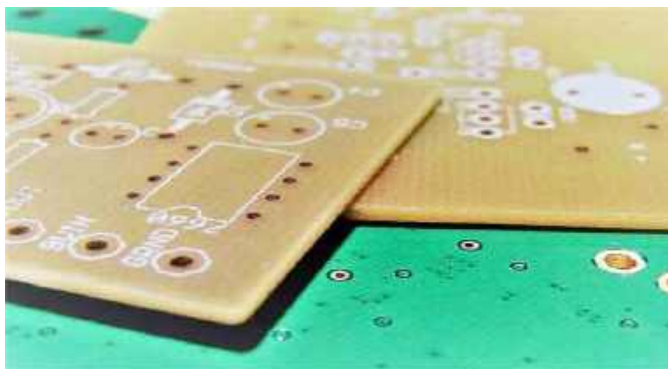
tel. +48 12 636 00 33

tel. +48 604 773 095

www.hatron.com

hatron@hatron.com





Fotografia 3. Płytki PCB wykonane z laminatu FR2 (na górze) i FR4 (na dole). Źródło: pcbmay.com

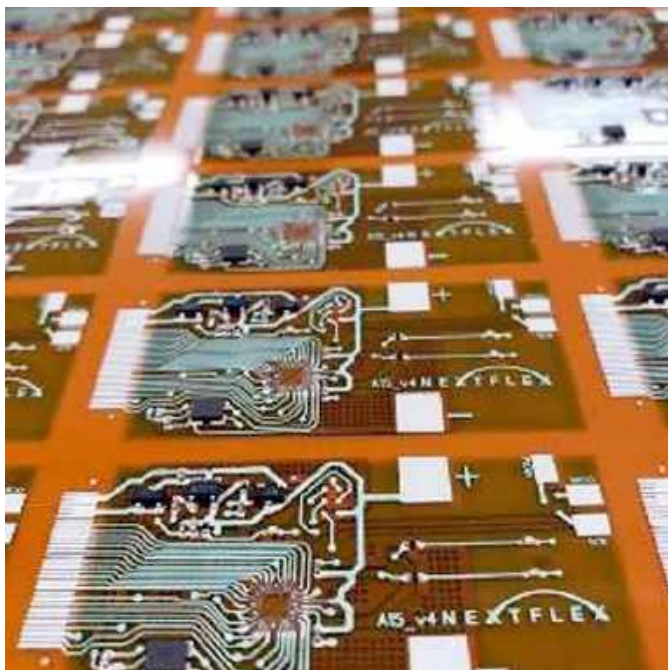
Również pliki frezowania dostarczane mogą być w formacie Excellon lub w formacie Gerber. Wszystkie opisane powyżej uwagi dotyczą także pliku frezowania.

Wybór materiałów

Produkcję płytki musimy zacząć od wyboru rodzaju laminatu, z jakiego zostanie ona wyprodukowana. Wybór ten skupia się głównie na decyzji co do materiału dielektrycznego, który oddziela od siebie poszczególne warstwy. Materiał ten pełni również funkcję podstawy mechanicznej – zapewnia odpowiednią sztywność PCB, a także odpowiada w pewnym stopniu za odprowadzanie ciepła z elementów na płytce. Poniżej scharakteryzowano kilka najczęściej stosowanych materiałów stosowanych do produkcji PCB.

1. Laminat szklano-epoksydowy FR-4

FR-4 to najbardziej popularny i szeroko stosowany materiał dielektryczny w produkcji PCB. Jest to kompozyt z włókna szklanego nasączony żywicą epoksydową, który charakteryzuje się dobrą odpornością termiczną, mechaniczną oraz relatywnie niskim kosztem. Materiał ten, dzięki dodatkowi związków bromu, jest trudnopalny, co sprawia, że jest bezpieczny w aplikacjach przemysłowych i konsumenckich. Oznaczenie FR-x pochodzi właśnie od stopnia trudnopalności (Flame Retarded). Z uwagi na te własności FR-4 jest obecnie „domyślnym” materiałem, z jakiego produkuje się PCB w większości standardowych urządzeń.



Fotografia 4. Elastyczne PCB wykonane z zastosowaniem laminatu poliimidowego

2. Inne materiały oparte na żywicach

Jak wspomniano wcześniej, oprócz laminatów szklano-epoksydowych (FR-4), w przemyśle spotyka się inne materiały oznaczone klasą FR. Są to głównie:

- FR1 – laminat papierowo-fenolowy,
- FR2 – laminat bawełniano-fenolowy,
- FR3 – laminat będący połączeniem materiałów stosowanych w FR1 i FR2,
- FR4 – opisany powyżej laminat szklano-epoksydowy,
- FR5 – laminat szklano-epoksydowy o zwiększonej temperaturze zeszklenia, charakteryzujący się jeszcze lepszą odpornością na wysokie temperatury.

Oprócz materiałów z tej grupy na rynku dostępne są inne materiały, stanowiące modyfikacje laminatu FR4, takie jak RF-35 czy getek. Zostały one zoptymalizowane do zastosowań w systemach wysokiej częstotliwości poprzez modyfikację lub zmianę żywicy epoksydowej na inną. Dostępny jest też laminat FR408HR o obniżonej stratności, co ma znaczenie w przypadku układów RF. Dodatkowo dostępne są laminaty o zastosowaniach specjalnych, m.in. oparte na bismaleimidowej żywicy triazynowej (przeznaczone – z uwagi na dobre parametry mechaniczne i dielektryczne przy jednocześnie zwiększonej wytrzymałości na wysoką temperaturę – do zastosowań w przemyśle kosmicznym).

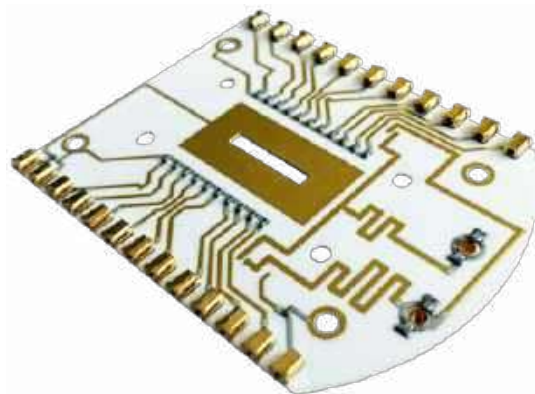
3. Poliimid (PI)

Poliimid (PI), znany również jako kapton, to materiał dielektryczny używany głównie do produkcji elastycznych PCB (typu flex i flex-rigid). Charakteryzuje się wyjątkową elastycznością, wysoką odpornością na temperaturę oraz doskonałymi właściwościami dielektrycznymi. Poliimid jest stosowany w miejscach, w których wymagane jest od PCB statyczne lub dynamiczne zginanie.

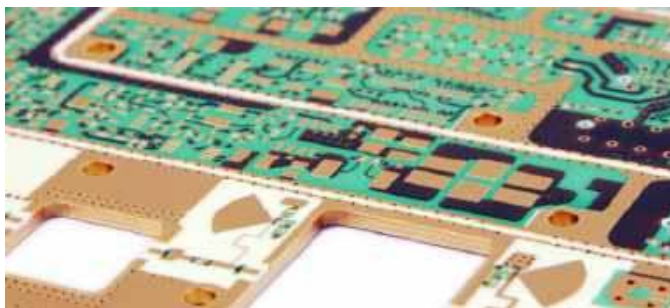
4. Ceramika

Płytki PCB wykonane na podłożach ceramicznych mają unikalne właściwości termiczne i elektryczne, co czyni je idealnym rozwiązaniem do zastosowań wysokotemperaturowych oraz w systemach wysokiej częstotliwości. Ceramika zapewnia doskonale odprowadzanie ciepła w urządzeniach, które generują dużą ilość energii – takich jak zasilacze, systemy radarowe czy moduły RF. Jednym z powszechnie używanych materiałów ceramicznych jest alumina lub tlenek glinu (Al_2O_3), ale rzadziej stosuje się azotek glinu (AlN), tlenek berylu (BeO), węgiel krzemu (SiC) oraz azotek boru (BN).

Płytki PCB wykonane na bazie materiałów ceramicznych charakteryzują się ekstremalną odpornością na wysoką temperaturę, dużą przewodnością cieplną i bardzo dobrymi parametrami elektrycznymi. Dodatkowo, z uwagi na fakt, że dielektryki w nich są jednorodne (podczas gdy np. w laminatach szklanych występują różnice w właściwościach dielektrycznych związane z periodyczną naturą tkanego włókna szklanego), idealnie nadają się do układów RF o ekstremalnie dużej gęstości połączeń.



Fotografia 5. Ceramiczna płytka drukowana (źródło: nextpcb.com)



Fotografia 6. Płytki drukowane na laminacie teflonowym (źródło: circuit-stampati.com)

Z drugiej strony, laminaty ceramiczne okazują się bardzo trudne w obróbce, a ich parametry mechaniczne są dokładnie takie, jakich można spodziewać się po ceramice. PCB tego typu są kruche i bardzo delikatne, a przy tym nie pozwalają na praktycznie żadne zginanie. Co za tym idzie, produkcja ceramicznych PCB jest wyjątkowo droga, a dodatkowo niektóre z materiałów ceramicznych są niebezpieczne dla ludzi – tlenek berylu (w sproszkowanej formie) ma działanie silnie rakotwórcze. Materiał w formie spieku w PCB nie jest groźny, pod warunkiem że nie jest obrabiany (wiercony, cięty) ani nie pękuje, gdyż wtedy też pyli(!).

5. PTFE (teflon – politetrafluoroetylen)

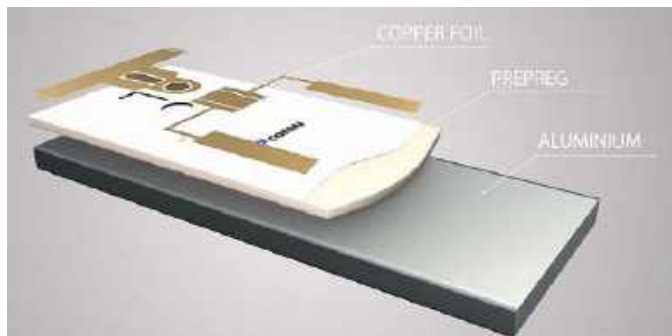
Teflon jest syntetycznym materiałem dielektrycznym stosowanym w PCB do systemów o wysokiej częstotliwości, szczególnie tam, gdzie wymagana jest minimalna strata sygnału oraz wysoka stabilność elektryczna. PTFE ma niski współczynnik stratności dielektrycznej, co sprawia, że jest idealny do aplikacji RF, systemów mikrofalowych czy telekomunikacyjnych. Do teflonu – na potrzeby produkcji podłoży do płytek – często dodaje się wypełniacze w postaci pyłów ceramicznych czy szklanych, co zwiększa jego stałą dielektryczną, nie zwiększając istotnie stratności.

Jakkolwiek teflon jest trudniejszy w obróbce niż laminaty szkło-epoksydowe, to PCB z PTFE są tańsze niż ceramiczne, a dodatkowo teflon wykazuje lepsze własności mechaniczne. Z uwagi na ten aspekt jest to obecnie drugi najpopularniejszy rodzaj laminatu po FR4, w szczególności w zastosowaniach RF czy kosmicznych.

6. Płytki drukowane z metalowym rdzeniem (MCPCB)

Płytki MCPCB korzystają ze specjalnego rodzaju podłoża – zamiast jednego z tradycyjnych materiałów dielektrycznych jest to metal, najczęściej aluminium lub miedź, pokryty tylko cienką warstwą dielektryka. Metalowy rdzeń PCB służy do bardzo efektywnego odprowadzania ciepła, co jest szczególnie ważne w aplikacjach, w których generowana jest spora moc strat (takich jak diody LED i przetwornice dużej mocy).

Po wybraniu materiału warstwy dielektrycznej należy zdecydować o jej pokryciu. Laminaty pokrywane są niemal wyłącznie miedzią, a jednym z nielicznych parametrów projektowych jest jej



Fotografia 7. Płytki MCPCB w przekroju (źródło: Cofan PCB)

grubość. Typowe laminaty, dostępne komercyjnie, pokryte są folią miedzianą o grubości 35 mikrometrów – to tak zwana grubość „jednej uncji”, ponieważ taką grubość ma folia miedziana o powierzchni jednej stopy kwadratowej i o wadze jednej uncji. Taka grubość jest typowo stosowana na ścieżki zewnętrzne w PCB; wewnętrzne ścieżki w płytkach wielowarstwowych zazwyczaj mają grubość „połowy uncji”, czyli 17,5 mikrometra.

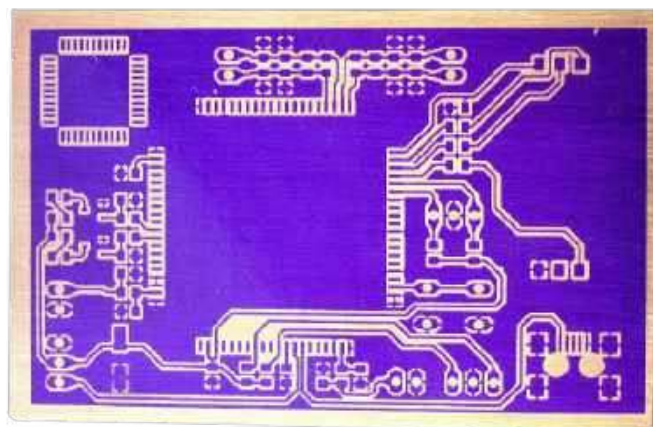
Norma IPC-2221 reguluje dobór pola przekroju miedzianych ścieżek do wymagań prądowych. W przypadku aplikacji o wysokiej gęstości połączeń, znacznej miniaturyzacji PCB czy po prostu w systemach dużej mocy, może nie być możliwe trasowanie dostatecznie szerokich ścieżek – w takim przypadku rekomenduje się zastosowanie płytek z grubszymi warstwami miedzi: 75 mikronów czy nawet więcej. W przemyśle spotyka się laminaty mające nawet do 1000 mikronów miedzi (1,0 mm!), nałożonych na warstwy dielektryczne. Produkcja PCB z takich substratów wymaga zastosowania specjalnych metod, np. zmodyfikowanego procesu trawienia lub dodatkowego platerowania.

Tworzenie obwodów

Na laminacie należy wytworzyć ścieżki. Istnieje kilka sposobów, z czego najpopularniejsze to trawienie (fotolitografia) lub frezowanie CNC.

Aby wzór ścieżek mógł być przeniesiony na PCB, na podstawie plików Gerber drukuje się maski do naświetlania płytek, pokrytych fotoczułym rezystem. Fotorezyst to polimer, który zmienia swoje właściwości chemiczne pod wpływem promieniowania ultrafioletowego (UV). Można go nakładać metodą natryskową lub w postaci cienkiego filmu, laminowanego na PCB. Istnieją dwa rodzaje fotorezystów: pozytywowe i negatywowe. W przypadku fotorezystu pozytywowego naświetlone obszary stają się rozpuszczalne, a w przypadku fotorezystu negatywowego – odporne na wytrawianie. W zależności od używanego rezystu, producent przygotowuje w odpowiedni sposób maski, tj. tak, aby w miejscach, które mają być trawione, fotorezyst stał się rozpuszczalny. Maski nakładane są na PCB i całość naświetlana jest promieniowaniem UV. Następnie naświetlone laminaty płukane są w rozpuszczalniku, który odsłania miejsca przeznaczone do usunięcia. Tak przygotowane laminaty gotowe są do trawienia.

PCB trawi się w specjalnym roztworze. Można wybierać pomiędzy dwoma rodzajami procesów trawiących – kwaśnymi lub zasadowymi. W przypadku kwaśnych stosuje się roztwory chlorku żelaza (FeCl_3) lub chlorku miedzi (CuCl_2), które często są dodatkowo podgrzewane. Istnieje szereg procesów, różniących się składem mieszanki (częstym dodatkiem jest chlorowódor (HCl), temperaturą trawienia (np. dla FeCl_3 optymalna temperatura procesu wynosi ok. 55°C) czy sposobem nanoszenia (laminaty mogą być zanurzane lub mieszanka trawiąca może być natryskiwana). Najpopularniejszy,



Fotografia 8. Płytki pokryta naświetlonym i wywołanym fotorezystem

w zastosowaniach przemysłowych, jest chlorek miedzi, gdyż zapewnia on dobrą prędkość trawienia, a jednocześnie pozwala na wierne oddanie drobnych detali rysunku ścieżek; jest to również jeden z tańszych procesów, który dodatkowo pozostawia po sobie relatywnie niewiele odpadów wymagających utylizacji.

W przypadku procesów alkalicznych (zasadowych) stosuje się mieszanek chorku miedzi, chlorowodoru i perhydrolu. Typowo wykorzystuje się tutaj trawienie natryskowe, gdyż pozwala to na lepszą kontrolę parametrów trawienia – zwłaszcza jego czasu, gdyż zbyt długie wystawienie PCB na roztwór trawiący może uszkadzać również warstwę dielektryczną. Proces ten jest bardzo szybki, jednak droższy niż w przypadku trawienia kwaśnego.

Oprócz metod chemicznych, do usuwania miedzi z miejsc niepokrytych fotorezystem używa się np. plazmy – jest to wydajny i bardzo precyzyjny proces, który dodatkowo zapewnia doskonałą geometrię (dzięki temu, że plazma trawi izotropowo), jednakże jest on bardzo drogi. Stosuje się go na ogół w sytuacjach, gdy zanieczyszczenia na PCB lub przetwarzanie odpadów procesowych stanowiłyby duży problem. Płytki przygotowane w ten sposób mają zerową szansę na zanieczyszczenie miedzi czy przelotek, co jest szczególnie istotne w przypadku systemów o wysokiej niezawodności (awionika, elektronika kosmiczna) i układów o ekstremalnie dużej gęstości połączeń.

Istnieją również metody mechaniczne, w których usuwa się niepotrzebną miedź z laminatu za pomocą np. frezowania CNC lub lasera. Tego rodzaju techniki są stosowane rzadziej, z uwagi na wyższy koszt i dużo gorsze skalowanie procesu, w porównaniu z fotolitografią. Dodatkowo przy frezowaniu CNC uzyskuje się na ogół gorszą rozdzielczość, nie nadaje się ono również do płytek z bardzo cienkimi ścieżkami. Tego rodzaju procesy stosowane są najczęściej do szybkiego wytwarzania płytek na potrzeby prototypu.

Kolejne warstwy

Po wytworzeniu ścieżek na warstwie miedzi, na płytkę z reguły nanosi się dwie kolejne warstwy – warstwę przeciwlutową (tzw. soldermaskę) oraz opis.

Soldermaska jest nakładana na PCB w celu ochrony ścieżek miedzianych przed korozją a także ułatwienia procesu lutowania. Proces zaczyna się od dokładnego oczyszczenia powierzchni PCB. Następnie na płytkę nanosi się cienką warstwę płynnej soldermaski, która może być aplikowana metodą sitodruku (szczególnie w przypadku produkcji na dużą skalę, w której opłacalne jest wykonanie stosownego sita), powlekania wałkiem lub metodą natryskową. Popularnym materiałem soldermaski jest fotopolimer epoksydowy, który utwardza się pod wpływem światła ultrafioletowego (UV).

Po nałożeniu soldermaski – o ile nie została ona naniesiona na PCB tylko w wybranych miejscach (sitodrukiem lub ew.

tampodrukiem) – na utworzoną warstwę nakłada się maskę fotolitograficzną, analogicznie jak w przypadku przygotowania laminatu do trawienia. Maskę wykonywana jest, oczywiście, na podstawie plików Gerber. W zależności od materiału soldermaski konieczne może okazać się zastosowanie maski pozytywowej lub negatywowej. Po naświetlaniu zmywa się nieutwardzony polimer z PCB za pomocą specjalnego roztworu, pozostawiając odsłonięte jedynie obszary, które mają być lutowane. Następnie soldermaska jest całkowicie utwardzana w piecu, co zapewnia jej zwiększoną trwałość i wytrzymałość.

Na kolejnym etapie na PCB nanosi się warstwę opisową, zwaną także silkscreenem (po angielsku oznacza to „sitodruk”, co powinno jednoznacznie wskazywać na proces, stosowany w druku tej warstwy). Historycznie rzecz biorąc, opis był właśnie nakładany sitodrukiem, jednakże obecnie coraz częściej zastępuje go druk wykonywany bezpośrednio na laminacie drukarkami atramentowymi, korzystającymi ze specjalnych tuszów akrylowych. Po nałożeniu farby na PCB warstwa opisowa jest utwardzana termicznie lub przy użyciu promieniowania UV, co zapewnia trwałość i odporność na działanie czynników zewnętrznych, takich jak ciepło czy chemikalia podczas lutowania.

Druk charakteryzuje się lepszą rozdzielczością i precyzją umiejscowienia (minimalna wielkość znaków wynosi 20 milów, zaś minimalna odległość pomiędzy liniami to 3 milsy) niż w przypadku sitodruku (odpowiednio: 35 milów i 5 milów). Dodatkowo druk atramentowy jest szybszy i tańszy, gdyż nie wymaga przygotowywania sit, co oznacza również, że jego koszt dobrze się skaluje, dzięki czemu często stosuje się go zwłaszcza w seriach prototypowych. Ponadto – z uwagi na wspomniane zalety – możliwe jest również zastosowanie tej metody druku do indywidualnego oznaczania PCB, na przykład poprzez nanoszenie unikalnych numerów seryjnych na każdej płytce.

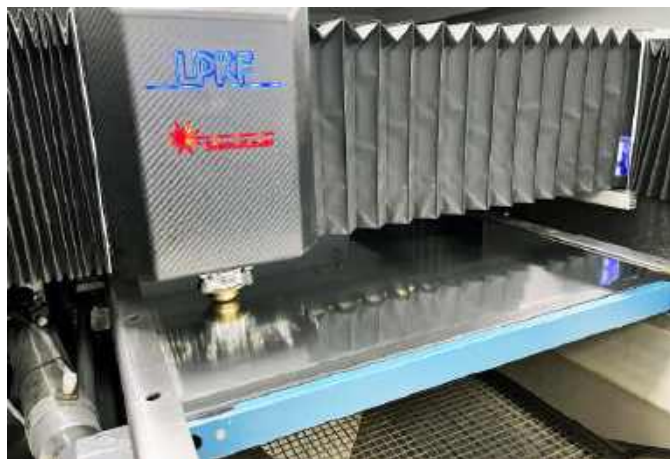
Wiercenie i metalizacja otworów

Wiercenie otworów pozwala na wykonanie zarówno przelotek, jak i otworów do montażu elementów przewlekanych czy przykręcenia PCB w obudowie. Proces ten realizowany jest za pomocą sterowanych komputerowo wiertarek lub frezarek. Maszyna wierci otwory na podstawie zdefiniowanego pliku owiertu (patrz opis wcześniej). W pliku tym, oprócz koordynatów każdego z otworów, znajduje się również informacja o jego średnicy. Obrabiarki najczęściej wykonują wszystkie otwory o jednej średnicy, po czym zmieniają narzędzie i rozpoczynają kolejny cykl wiercenia. Z tego też powodu najlepiej jest ograniczyć liczbę średnic wierconych otworów do bezwzględnie minimum.

Po wykonaniu otworów następuje proces metalizacji, polegający na pokryciu ścianek otworów cienką przewodzącą warstwą miedzi, która tworzy połączenie elektryczne. Proces ten rozpoczyna się od chemicznego osadzenia bardzo cienkiej warstwy miedzi w otworach, co umożliwi później elektrochemiczne naniesienie grubszej warstwy miedzi.

Omówione powyżej procesy stanowią podstawowe etapy produkcji płytek drukowanych. Producenci oferują również szereg procesów opcjonalnych. Niektóre z nich – jak cynowanie pól lutowniczych – są stosowane bardzo często, a inne – jak np. zasklepianie przelotek w polach lutowniczych – są stosowane tylko w określonych przypadkach.

Wspomniane powyżej cynowanie powierzchni miedzi jest niemalże domyślną opcją pokrycia. Warstwa cyny na miedzi zapewnia jej ochronę przed utlenianiem, a także poprawia lutowalność. Istnieją dwa podstawowe sposoby cynowania PCB – chemiczne oraz termiczne za pomocą gorącego powietrza (HAL). Pierwszą wspomnianą metodę stosuje się do delikatniejszych płytek, jednakże wytwarzana warstwa cyny jest cieńsza. Warstwa ta z czasem się utlenia – dlatego pokryte nią płytki nadają się do użytku (szczególnie przy



Fotografia 9. Płoter laserowy do wycinania szablonu do nakładania maski przeciwlutowej (źródło: Semicon)

lutowaniu maszynowym) tylko przez kilka miesięcy, jeśli przechowywane są w powietrzu atmosferycznym. Cynowanie HAL zapewnia stabilną warstwę miedzi zabezpieczającą PCB, która zachowuje lutowalność dużo dłużej.

Innym pokryciem miedzi na PCB, stosowanym w przemyśle, jest złocenie. Najczęściej złoci się elementy stykowe płytki – złącza krawędziowe, przyciski (jeśli wykonywane są na PCB) itp. Ma to zapewnić dłuższą żywotność tych elementów, zwiększając ochronę przed wpływem środowiska i poprawiając przewodność. Złocenie takie wykonuje się elektrochemicznie, może być ono selektywne.

Do bardziej złożonych procesów należy np. zatykanie przelotek. Jest to konieczne, jeśli przelotka znajduje się na polu lutowniczym elementów SMD – inaczej spowodować może niekontrolowany wpływ lutownia na drugą stronę PCB, co może pogorszyć własności spoiny lub (w skrajnych wypadkach) spowodować, że w ogóle nie zapewni ona połączenia elektrycznego. Przelotki tego typu najczęściej występują np. pod układami BGA lub innymi elementami o wysokiej gęstości wyprowadzeń. Aby przygotować takie pole lutownicze, przelotki po metalizacji zapelnia się żywicą epoksydową, a następnie ponownie metalizuje powierzchnię padów – dzięki czemu żywica nie pozostaje odsłonięta, a pole lutownicze jest jednolite.

Korzystanie z usług EMS – na co zwrócić uwagę?

W przypadku produkcji urządzeń elektronicznych (niezależnie od tego, czy chodzi o prototypy, czy masowe serie) często korzystamy z usług producentów kontraktowych (EMS, ang. Electronics Manufacturing Services – usługi produkcji elektronicznej), czyli firm, które zajmują się produkcją PCB na zlecenie. Na rynku istnieje bardzo wiele tego rodzaju przedsiębiorstw – krajowych i zagranicznych. Różnią się one poziomem cen, zakresem oferowanych usług itp., jak jednak wybrać tego jedyne go dostawcę, któremu powierzymy produkcję naszego projektu?

Wybór dostawcy EMS

Wybór EMS ma bezpośredni wpływ na jakość i niezawodność gotowych produktów elektronicznych. Jednym z pierwszych aspektów, na które należy zwrócić uwagę, jest doświadczenie dostawcy. Ważne, aby firma miała udokumentowane doświadczenie w branży elektronicznej, najlepiej poparte szerokim portfolio dotychczasowych realizacji. Warto sprawdzić również polecenia od innych klientów z branży, szczególnie tych działających w podobnym sektorze co my lub mających podobne wymagania technologiczne. Dobrze jest ocenić, czy dostawca miał doświadczenie w realizacji projektów o różnym stopniu skomplikowania, zwracając szczególnie uwagę na zamówienia o podobnym do naszego poziomie zaawansowania i złożoności. Dodatkowo – jeśli działamy w sektorze, który nakłada na producentów pewne dodatkowe wymagania, takim jak sektor urządzeń medycznych czy motoryzacyjnych – warto jest się upewnić, czy EMS posiada w nim doświadczenie.

Kolejnym istotnym czynnikiem okazuje się zgodność z normami i certyfikaty, które posiada EMS. W branży elektronicznej kluczowe są certyfikaty jakości, takie jak ISO 9001 (potwierdzające wdrożenie standardów zarządzania jakością), a także np. ISO 13485 (obejmujący wymagania dla firm produkujących urządzenia medyczne). Ponadto istotne są certyfikaty IPC, dotyczące standardów produkcji elektroniki, a krytyczna jest zgodność z dyrektywą RoHS, regulującą stosowanie niebezpiecznych substancji (między innymi ołowiu w stopach lutowniczych) w produkcji elektroniki. W przypadku aplikacji lotniczych czy motoryzacyjnych warto także zwrócić uwagę na normy AS9100 czy IATF 16949, obowiązujące w tych sektorach. Bez spełniania wspomnianych norm przez EMS wprowadzenie naszego urządzenia do sprzedaży czy finalna jego certyfikacja mogą być poważnie utrudnione lub wręcz niemożliwe.

Równie ważne są możliwości produkcyjne dostawcy. Warto zwerifikować, czy firma jest w stanie sprostać zarówno produkcji

prototypowej, jak i wielkoseryjnej. Dostawca EMS powinien dysponować nowoczesnymi liniami montażowymi, które pozwalają na montaż powierzchniowy (SMD), jak i przewlekany (THT), dzięki czemu możliwa jest realizacja różnorodnych projektów.

Dobrze, jeśli dostawca oferuje także dodatkowe usługi (w tym: testowanie obwodów, zarządzanie łańcuchem dostaw, montaż końcowy czy usługi związane z magazynowaniem komponentów), jeśli takowych potrzebujemy. Uprości to przejście od fazy prototypowej do masowej produkcji, a także usprawni elastyczną produkcję masową.

Dodatkowo ważnym aspektem, który może decydować o wyborze dostawcy, jest jego zdolność do zarządzania ryzykiem w zakresie dostaw komponentów. Dostawcy EMS często oferują wsparcie w zakresie pozyskiwania komponentów elektronicznych, monitorowania ich dostępności, a także zarządzania wycofywaniem produktów (EOL – End of Life). W przypadku branż wymagających szybkich reakcji, m.in. w branży elektroniki konsumenckiej, dostawca musi mieć odpowiednie zasoby i elastyczność, aby reagować na dynamiczne zmiany w dostępności komponentów, zapewniając ciągłość produkcji.

Kwestie techniczne i logistyczne

Design for Manufacturing (DFM) to paradygmat, który polega na optymalizacji projektu pod kątem łatwości i efektywności produkcji urządzeń, co ma ogromne znaczenie na każdym etapie cyklu życia produktu. Konsultacje z dostawcą usług EMS na wczesnych etapach projektowania są niezwykle istotne, ponieważ pozwalają uniknąć błędów, które mogą znacznie podnieść koszty produkcji lub opóźnić realizację projektu. Specjaliści z EMS mogą doradzić, jak zoptymalizować rozmieszczenie komponentów na płytce, jakie technologie montażu wybrać (np. SMD czy THT), a także jakie standardy jakościowe warto uwzględnić. Dzięki temu projekt nie tylko spełnia oczekiwania funkcjonalne, ale również jest łatwy i tani do wdrożenia w produkcji seryjnej. Właściwe podejście DFM pomaga również zmniejszyć liczbę niezgodności, zredukować ilość odpadów powstających przy produkcji oraz zwiększać ogólną skuteczność procesu produkcji, co przekłada się finalnie na niższe koszty i krótszy czas realizacji.

Minimalizacja ryzyka związanego z dostępnością komponentów to kolejny istotny element współpracy z dostawcą EMS. Branża elektroniki jest dynamiczna, a wiele komponentów może z dnia na dzień stać się trudno dostępnych, zostać wycofanych z rynku (End of Life, EOL) lub mieć długi czas oczekiwania na dostawę. Wszyscy mieliśmy okazję zaobserwować to podczas pandemii w 2020 roku. Odpowiedni partner EMS może pomóc nam znaleźć alternatywne rozwiązania, sugerując zamienniki komponentów, które są dostępne na rynku lub mają krótszy lead time. Dodatkowo, dzięki szerokim kontaktom i relacjom z dostawcami komponentów, firmy EMS mogą również często negocjować dla nas lepsze warunki zakupu czy też szybciej reagować na problemy związane z łańcuchem dostaw. Minimalizacja ryzyka obejmuje również wsparcie w zakresie prognozowania popytu i planowania zapasów komponentów w celu uniknięcia opóźnień produkcyjnych i zapewnienia ciągłości dostaw.

Czas realizacji, czyli lead time, pozostaje niezwykle istotnym parametrem wpływającym na terminowość dostarczenia gotowego produktu do klienta. Analiza lead time obejmuje nie tylko czas potrzebny na produkcję, ale także na projektowanie, pozyskanie komponentów, montaż, testowanie układów, integrację całego urządzenia i dostawę go do klienta. Każdy z tych etapów może wpłynąć na ostateczny czas realizacji, dlatego tak ważne jest ścisłe zarządzanie całym procesem. Współpraca z EMS może istotnie pomóc w optymalizacji niemalże każdego z tych kroków – od skrócenia czasu prototypowania, poprzez zautomatyzowanie procesów montażowych, aż po wdrożenie nowoczesnych narzędzi do monitorowania

produkcji w czasie rzeczywistym. Skrócenie lead time często zależy od efektywnego zarządzania łańcuchem dostaw, a także od zdolności dostawcy EMS do szybkiego skalowania produkcji, gdy jest to wymagane. Z tego powodu firmy projektujące produkty elektroniczne muszą ściśle współpracować z wybranym EMS, aby zapewnić korzystne terminy dostaw nowych produktów na rynek i zminimalizować ryzyko ewentualnych opóźnień.

Dodatkowe usługi oferowane przez EMS

Oprócz produkcji samych płytek drukowanych, firmy EMS oferują szeroką gamę usług. Podstawowe z nich są związane z produkcją elektroniki – to m.in. montaż elementów, testowanie gotowych modułów, programowanie firmware, ale także możliwość montażu elementów mechanicznych, integracja całego urządzenia czy nawet jego konfekcjonowanie i/lub pakowanie. Przyjrzyjmy się bliżej, jakie są to usługi, czym się charakteryzują i na co należy zwrócić uwagę.

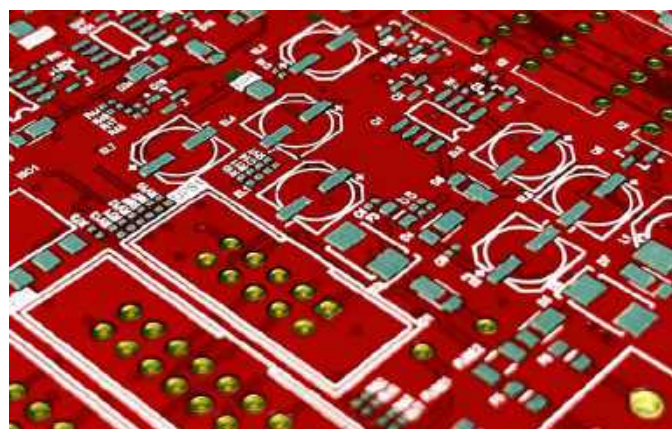
Montaż elementów na PCB

Montaż elementów na płycie to podstawowa usługa – oprócz produkcji PCB – oferowana przez firmy EMS. Jest to bardzo szerokie zagadnienie, wykraczające poza objętość tego artykułu, dlatego procesowi montażu i lutowania elementów na PCB przyjrzymy się w pewnym uproszczeniu. Proces ten obejmuje zarówno montaż komponentów powierzchniowych (SMD), jak i przewlekanych (THT), jednakże oba te procesy istotnie się od siebie różnią.

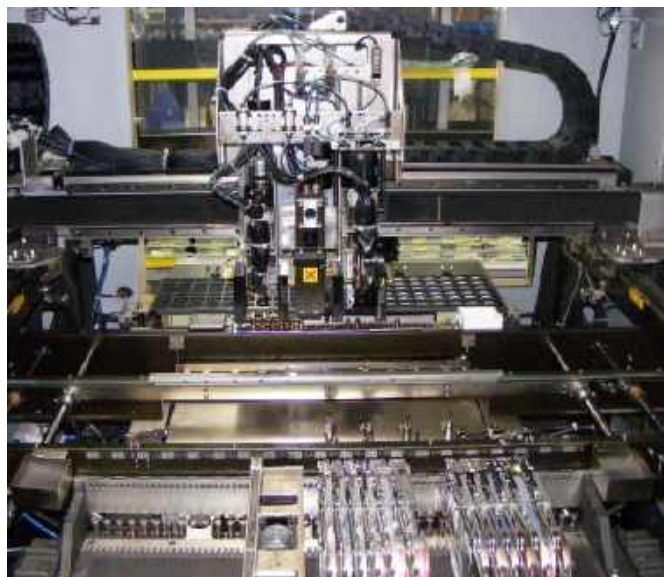
Montaż elementów SMD odbywa się w dużej mierze automatycznie. Proces ten wymaga kilku kroków.

1. Nakładanie pasty lutowniczej

Pierwszym krokiem jest przygotowanie PCB do lutowania poprzez naniesienie pasty lutowniczej. Pastę nanosi się metodą podobną do sitodruku. Na płycie PCB umieszcza się metalową maskę, zwaną szablonem – ma ona otwory dokładnie odpowiadające miejscom, w których zostaną umieszczone pady pod komponenty. Maszyna nakłada pastę lutowniczą równomiernie na powierzchnię szablonu. Pasta penetruje otwory w szablonie, osadzając się jedynie na padach PCB, a jej nadmiar ściągany jest przez maszynę z powierzchni szablonu. Po usunięciu szablonu pasta pozostaje na wyznaczonych miejscach, gotowa do przyjęcia elementów SMD, które następnie będą lutowane w procesie rozpliwowym. Sam szablon jest wcześniej wycinany na podstawie odpowiedniego pliku Gerber. Grubość blachy specyfikujemy przy składaniu zamówienia, ponieważ to głównie od niej zależy ilość pasty nakładanej na pola lutownicze (parametr ten dobiera się odpowiednio do rodzaju i gęstości upakowania elementów na PCB, a także ich rastra). Szablony produkuje się z blach ze stali nierdzewnej o grubości od 0,05 mm do 0,30 mm; typowo używa się szablono



Fotografia 10. PCB po nałożeniu pasty lutowniczej przez szablon (źródło: Semicon)



Fotografia 11. Maszyna pick-and-place do układania elementów na PCB (źródło: Wikimedia.org)

2. Nakładanie kleju

Czasami elementy SMD dodatkowo przykleja się do PCB przed lutowaniem. Postępuje się tak z cięższymi elementami, aby nie narażać połączenia lutowanego na nadmierne obciążenia mechaniczne. Dodatkowo, jeśli elementy SMD znajdują się po obu stronach płytki, klei się je, aby nie odpadały, gdy będą się znajdowały na dolnej stronie PCB.

Klej nakłada się za pomocą maszyny dozującej sterowanej numerycznie. Precyzyjnie aplikuje ona kroplę kleju w wyznaczonych miejscach na płycie drukowanej. Miejsca, w których naniesiony ma być klej, definiowane są – analogicznie, jak w przypadku pasty itp – w odpowiednim pliku Gerbera. Klej zostaje później utwardzony w piecu, co zapewnia trwałe przymocowanie elementów do PCB.

3. Układanie elementów

Na tak przygotowaną płytkę nakłada się elementy. W tym celu stosuje się tak zwane maszyny pick-and-place, które z niezwykłą precyzją umieszczają komponenty na odpowiednich miejscach.

Konfiguracja systemu układania elementów zapisywana jest w postaci dwóch plików – wykazu elementów (BoM – Bill of Materials) oraz pliku konfiguracyjnego pick-and-place. W wykazie znajduje się spis wszystkich elementów: dla każdego z nich opisana musi być co najmniej nazwa – oznaczenie identyfikacyjne, które pozwala na znalezienie go na płycie (np. R1, R2, Q10, IC67, etc.) oraz oznaczenie, najlepiej w postaci numeru produktu, identyfikującego konkretny element konkretnego producenta. Nazwa jest potrzebna, gdyż w drugim pliku znajdują się nazwy oraz koordynaty na PCB i kąt obrotu elementu.

Maszyna pick-and-place pobiera – z rolki lub z tacki – odpowiedni element za pomocą chwytaka pneumatycznego, a następnie obraca go, aby znalazł się w odpowiedniej orientacji. Na tym etapie często korzysta się z systemu wizyjnego do kontroli ustawienia elementu w uchwycie przed umieszczeniem go na PCB. Po potwierdzeniu, że element znajduje się poprawnie w uchwycie, maszyna przemieszcza uchwyt w zadane koordynaty i umieszcza delikatnie element na powierzchni PCB.

Jakkolwiek typowe maszyny pick-and-place są bardzo szybkie i są w stanie ułożyć od stu do nawet ponad tysiąca elementów w minutę, to czas realizacji tego kroku mocno zależy od projektu PCB. Jeśli chcemy optymalizować płytkę pod kątem czasu produkcji, warto rozważyć takie ustawienie elementów, które ułatwia pracę maszyny pick-and-place. Podstawowym zabiegiem jest układanie ich w taki sposób, aby zminimalizować konieczność



Fotografia 12. Piec do lutowania rozpliwowego (źródło: Semicon)

obracania. Normy JEDEC definiują, jaka orientacja jest określana jako 0° dla każdej obudowy SMD.

4. Lutowanie w piecu

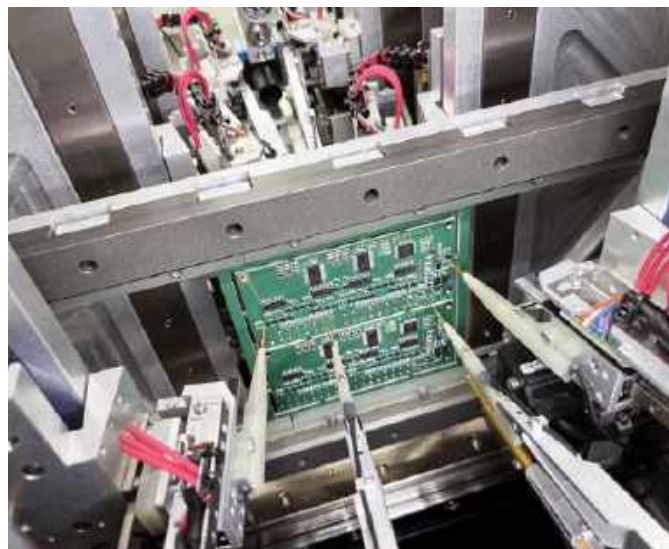
Gdy elementy znajdują się na swoich miejscach, płytka przechodzi przez piec lutowniczy, gdzie w kontrolowanych warunkach odbywa się lutowanie. Proces ten nosi nazwę lutowania rozpliwowego (reflow soldering) i zapewnia stabilne oraz trwałe połączenie elektryczne między komponentami a padami na płytce. Sposób lutowania – temperatura, czasy przebywania w poszczególnych strefach etc. – definiowane są na ogół w zależności od używanego lutownia. Opis tych zależności znaleźć można w normach IPC-7530A oraz IPC/JEDEC J-STD-020.

Z kolei montaż elementów THT, czyli przewlekanych, wymaga ręcznego lub półautomatycznego umieszczenia komponentów w otworach przelotowych płyty. Elementy te są następnie lutowane za pomocą tzw. procesu lutowania na fali (wave soldering), w którym cała płytka przepływa nad falą roztopionego lutownia. Ta metoda zapewnia równomierne i mocne połączenia mechaniczne oraz elektryczne, szczególnie w przypadku większych i bardziej wytrzymałych komponentów, jak złącza, kondensatory elektrolityczne czy transformatory. Montaż THT jest jednak bardziej czasochłonny i pracochłonny (w porównaniu z technologią SMT), dlatego stosuje się go głównie tam, gdzie wymagana jest większa wytrzymałość mechaniczna połączeń lub w sytuacjach, w których (z racji charakterystyki komponentów) nie da się zastosować montażu powierzchniowego. Pomimo że montaż przewlekany wydaje się starszą technologią, to nadal w wielu przypadkach okazuje się on niezastąpiony, a obie techniki często współistnieją w ramach jednej płytki PCBA.

Testowanie i kontrola jakości

Testowanie PCB i PCBA to kluczowy etap produkcji, który ma na celu weryfikację poprawności wykonania płytki (PCB) oraz płytki obsadzonej elementami (PCBA). Stosowane są tu różne metody, które pomagają upewnić się, że gotowy produkt działa zgodnie z założeniami.

Jedną z najczęściej używanych metod jest **Automatyczna Inspekcja Optyczna** (AOI), która polega na zastosowaniu kamer wysokiej rozdzielczości, obrazujących PCB lub PCBA i w celu porównania zarejestrowanego obrazu z wzorcem. AOI jest w stanie wykryć błędy, np. w postaci brakujących komponentów, uszkodzeń ścieżek, złej orientacji elementów, a także niektóre wady lutowania. Tego rodzaju inspekcja jest bardzo szybka i skuteczna w wykrywaniu defektów, które są widoczne na PCB, dlatego jest szeroko stosowana na różnych etapach produkcji.



Fotografia 13. Tester z sondą latającą (źródło: Semicon)

Kolejną popularną metodą stanowi **testowanie elektryczne** (In-Circuit Testing – ICT). W testerach ICT do gotowej płytki przykładają się specjalne sondy, które mierzą poszczególne połączenia i sprawdzają, czy obwody na PCB są poprawnie wykonane, czy nie ma pomiędzy nimi przerw i zwarców. W przypadku badania PCBA można również określić, czy komponenty działają prawidłowo. ICT może być wykonywane na specjalnych platformach testowych, które mają dostęp do każdego wymaganego pola lutowniczego i połączenia, jednak procedura taka wymaga stworzenia specjalnych narzędzi (tzw. nailbed – ang. „łóżka z kolcami”, ze względu na wizualne podobieństwo tego rodzaju matrycy do łóżka fakira), co sprawia, że opisane podejście do ICT jest opłacalne głównie przy produkcji masowej.

Alternatywę dla zastosowania nailbedów stanowi **testowanie z sondą latającą**. Nie wymaga ono specjalnego uchwytu – zamiast niego używa się ruchomych sond, które wędrują po płytce i testują PCB lub PCBA poprzez przyłożenie próbników do poszczególnych punktów testowych. Ta metoda jest bardziej elastyczna i tańsza w przypadku mniejszych serii, ale jest też wolniejsza, niż testowanie modułu na nailbedzie. Flying probe bywa często stosowana w odniesieniu do prototypów i produkcji niskonakładowych.

REKLAMA

OBWODY DRUKOWANE
Produkcja, Projektowanie, Montaż

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 czelone FR4

Serie dowolne
Prototypy
Maksymalny wymiar płytek 1w 630 mm

Dokumentacja technologiczna
Dokumentacja konstrukcyjna
Trawione szablony SMD

Montaż elektroniczny
Krótkie terminy
Wykonania super ekspresowe

Aktywny kalkulator prototypów na stronie internetowej

Pokrycie Sn lub SnPb inne na życzenie
Maski, opisy montażowe w różnych kolorach

Płytki drukowane są przedmiotem wielu norm. Najpopularniejsze z nich to standardy opracowane przez IPC (Institute for Printed Circuits), w tym:

- IPC-A-600 – norma opisująca akceptowalność płytek drukowanych, dotycząca wizualnych i mechanicznych cech płytki.
- IPC-A-610 – norma do oceny akceptowalności montażu komponentów na PCB. Określa kryteria dotyczące jakości lutowania, montażu elementów, a także typowe wady ich montażu.
- IPC-6012 – specyfikacja dotycząca wymagań dla sztywnych PCB, obejmująca kwestie mechaniczne, elektryczne i środowiskowe.
- IPC-7711/7721 – norma dotycząca naprawy i modyfikacji płytek oraz elementów.

Spełnienie tych norm zapewnia, że produkty są zgodne ze światowymi standardami jakości, co jest niezwykle istotne w branżach o wysokich wymaganiach, takich jak elektronika medyczna, motoryzacyjna czy wojskowa (pamiętajmy jednak, że niektóre z tych sektorów mogą mieć swoje własne, dodatkowe wymagania co do urządzeń elektronicznych).

Konfekcjonowanie i integracja

Usługi integracji końcowej obejmują czynności związane z ostatecznym montażem i przygotowaniem produktu do użytkowania. Etap ten może obejmować montaż PCBA w obudowie, instalację elementów mechanicznych, złączy, paneli sterujących czy wyświetlaczy, a także finalne testowanie gotowych produktów. Często jednym z ostatnich kroków jest wgranie oprogramowania/firmware do pamięci znajdujących się na PCBA mikrokontrolerów itp. Proces ten odbywa się przy zastosowaniu specjalistycznych narzędzi (programatorów) i oprogramowania, zależnie od tego, jakie układy programowalne znajdują się w systemie. Najczęściej programowanie systemu stanowi część testów elektrycznych systemu.

Konfekcjonowanie produktów w branży EMS obejmuje szeroki zakres działań związanych z przygotowaniem gotowych wyrobów do transportu i sprzedaży. Jednym z niepomijalnych aspektów jest pakowanie PCB oraz PCBA w opakowania antystatyczne – ze względu na potrzebną ochronę wrażliwych

komponentów elektronicznych przed destrukcyjnym wpływem ładunków elektrostatycznych. Używane są w tym celu np. specjalne torebki ESD oraz pianki zabezpieczające. W przypadku produktów, które wymagają dodatkowej ochrony przed wilgocią, stosuje się hermetyzację, czyli zamknięcie ich w szczelnych opakowaniach, często z pochłaniaczami wilgoci. Tego typu rozwiązania są niezbędne, zwłaszcza przy produktach wysyłanych do odległych lokalizacji lub składowanych w trudnych warunkach środowiskowych.

Dzięki kompleksowym usługom konfekcjonowania i integracji klienci otrzymują gotowe do użytku produkty prosto z zakładu EMS, co minimalizuje konieczność angażowania zamawiającego w dodatkowe etapy produkcji i logistykę. Jest to szczególnie istotne dla mniejszych firm, dla których uruchomienie działu produkcyjnego byłoby najwyczejniej mniej opłacalne.

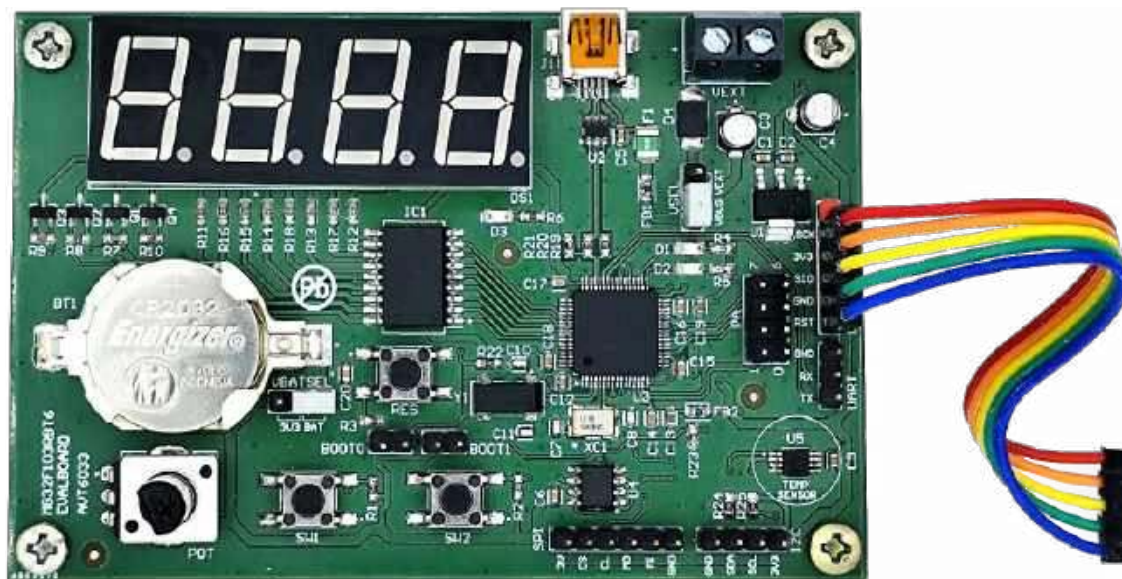
Podsumowanie

Produkcja płytek drukowanych i urządzeń elektronicznych to złożony proces, który obejmuje wiele etapów – od projektowania PCB, przez montaż komponentów, aż po testowanie i integrację gotowego produktu. Współczesne płytki drukowane są sercem praktycznie wszystkich urządzeń elektronicznych; ich produkcja wymaga zaawansowanego sprzętu oraz rygorystycznych standardów jakości. Dzięki nowoczesnym technologiom, takim jak automatyczny montaż SMD, zautomatyzowane inspekcje – optyczne czy elektryczne – możliwe jest tworzenie skomplikowanych układów elektronicznych, przy zachowaniu wysokiego współczynnika sprawnych produktów.

Współpraca z firmami EMS pozwala na optymalizację każdego etapu produkcji, oferując nie tylko wsparcie w projektowaniu, ale również zarządzanie całym łańcuchem dostaw, programowanie firmware'u czy nawet konfekcjonowanie produktów. Ostatecznie, kluczowe dla sukcesu jest zapewnienie wysokiej jakości i zgodności z normami (np. IPC), które gwarantują, że finalny produkt spełnia standardy niezawodności oraz bezpieczeństwa, niezależnie od skali produkcji.

Nikodem Czechowski

REKLAMA



Kurs programowania mikrokontrolerów Megawin

Na łamach „Elektroniki Praktycznej” publikujemy pierwszy na świecie kurs podstaw programowania interesujących, budżetowych mikrokontrolerów z rdzeniem ARM Cortex-M3 firmy Megawin.



ulubionykiosk.pl

Udane prototypy sposobem na osiągnięcie celów w zaawansowanych projektach



WÜRTH
ELEKTRONIK
MORE THAN
YOU EXPECT

Wyobraźmy sobie następującą sytuację: termin wykonania złożonej płytki PCB został dotrzymany, prototypy dostarczono i zmontowano terminowo. Teraz czas na testy środowiskowe: czy projekt spełni wszystkie wymagania, czy też niezbędne okaże się wprowadzenie poprawek? Wystarczy niewielka optymalizacja czy może konieczny będzie redesign? Czy rozwiązanie można bez problemu wdrożyć do produkcji seryjnej i obsłużyć już istniejące zamówienia klientów?



Im bardziej złożona technologia, tym większy wpływ na termin dostarczenia gotowego produktu mają wszelkie opóźnienia, ponieważ w proces zaangażowanych jest wiele działów i partnerów. Wczesne rozpoczęcie współpracy z doświadczonym i wydajnym producentem płytek PCB jest kluczowe, aby już na początkowym etapie rozpoznać i przezwyciężyć ewentualne problemy.

6 kroków do sukcesu:

1. Koncepcja. Doradztwo techniczne, symulacje, próbki PCB oraz webinaria wspierają proces wyboru technologii. Techniczne zarządzanie projektami i silny, lokalny zespół sprzedaży zapewniają niezbędną wiedzę specjalistyczną. Pomagamy naszym klientom zoptymalizować ich projekty pod kątem wytwarzania (DFM).

2. Rozwój i specyfikacja. Würth Elektronik oferuje różne technologie – od BASIC, HDI i RIGID.flex do HIGH.speed. Posiadamy certyfikaty zgodne z najwyższymi standardami branżowymi: EN 9100 i IATF 16949. Już na wczesnym etapie projektu zapewniamy wsparcie naszym doświadczeniem w zakresie procedur testowych, takich jak „interconnect stress test”, aby pomóc w osiągnięciu zakładanych wyników.

3. Projekt PCB. Standardy zapewniają wysoką jakość, a zarazem obniżają koszty. Skorzystaj z naszych zasad projektowania i gotowych stackupów. Nasi specjaliści techniczni oferują indywidualne rozwiązania pod specjalne wymagania.

4. Prototypy. Ekspresowe dostawy – już od trzech dni – ze sklepu WEdirekt (www.wedirekt.com) umożliwiają elastyczne reagowanie na dynamicznie zmieniającą się sytuację i pozwalają na przyspieszenie wstępnych prac nad projektem. Liczbę zapytań i wyjaśnień dotyczących danych produkcyjnych można zminimalizować (lub – w idealnym przypadku – w ogóle uniknąć tego etapu komunikacji) poprzez ścisłą koordynację na etapach 2 i 3.

5. Zatwierdzenie próbek i kwalifikacja produktu. Dostosowane do potrzeb testy i raporty, np. zgodne z klasą 3 wg IPC-A-600, zapewniają niezawodność konstrukcji.

6. Seria. Dzięki dekadom zbierania doświadczeń Würth Elektronik oferuje niezawodną produkcję w Niemczech i Azji – od pierwszej serii, przez kolejne rewizje, aż do końca cyklu życia produktu.

Dzięki temu projekty są realizowane terminowo, z zapewnieniem najwyższej jakości produktów.

Zainteresowany? Skontaktuj się z naszymi przedstawicielami:

Tomasz Renkiel – Area Sales Manager Southern Poland
+48 785 085 700, Tomasz.Renkiel@we-online.com

Marcin Ziolkowski – Area Sales Manager Northern Poland
+48 693 910 475, Marcin.Ziolkowski@we-online.com

Informacje dotyczące naszego portfolio produktowego: www.we-online.com/pcb

Próbki PCB, wykonane w naszych flagowych technologiach, dostępne są pod tym linkiem: www.we-online.com/physical-pcb-samples

REKLAMA

KOMPETENCJA W PRODUKCJI PCB OD 1971

Od standardowych laminatów wielowarstwowych do technologii high-end!



WÜRTH
ELEKTRONIK
MORE THAN
YOU EXPECT

EN 9100 certified

www.we-online.com

Internet Rzeczy w pomiarach środowiskowych (10)

Czujnik UV z fotodiodą GUYA-S12SD

Określanie poziomu promieniowania UV może być istotne przy słonecznej pogodzie. Pomagają w nim analogowe i cyfrowe czujniki UV. Fotodiodę GUYA-S12SD, czułą w pasmie nadfioletu, zaimplementowano w wielu modułach dostarczanych przez różne firmy. Jednym z ciekawszych jest moduł „UV Sensor” firmy Waveshare. Został on dołączony do płytki Enviro Weather (czujniki: ciśnienie, temperatura i wilgotność z BME280; poziom oświetlenia z LTR-559) wraz z modułem Grove – I²C Color Sensor (czujnik RGB TCS34725) i modułem AS7341 „Spectral Color Sensor” (czujnik wielospektralny AS7341). Do pokazywania rezultatów pomiarów służy dosyć spory (i niedrogi) wyświetlacz e-Paper.

Czujniki UV

Scalone czujniki nadfioletu są produkowane przez różne firmy, w tym: ams-OSRAM, Hamamatsu, Liteon, Silicon Labs czy Genicom. Układy te dostarczają pomiar w postaci analogowego napięcia lub cyfrowo – poprzez szynę I²C. Jednak zarówno same układy scalone, jak i moduły uruchomieniowe bywają często wycofywane, niedostępne lub bardzo drogie. W zasadzie na rynku da się zdobyć moduły z trzema czujnikami: analogowym czujnikiem GUYA-S12SD firmy Genicom (UVA: 240...370 nm), cyfrowym sensorem LTR-390 firmy Liteon (UVA+UVB: 280...430 nm) oraz trzykanałowym czujnikiem AS7331 firmy ams-OSRAM (UVA, UVB, UVC).

Jednokanałowy czujnik LTR-390 firmy Liteon pozwala na pomiar natężenia promieniowania ultrafioletowego i światła widzialnego (ALS) – cechuje się on wysoką czułością, dużym zakresem odpowiedzi liniowej (1:18 000 000), szybką reakcją oraz małym rozmiarem. Układ ma wbudowany przetwornik ADC z maksymalną rozdzielczością 20 bitów. Moduł jest zasilany napięciem od 3,3 V do 5 V, a komunikacja odbywa się przez interfejs I²C. Układ nie wymaga kalibracji.

Trzykanałowy czujnik AS7331 firmy ams-OSRAM to zintegrowany detektor UV o niewielkim poborze mocy i niskim poziomie szumów [4]. Trzy oddzielne kanały UVA, UVB i UVC konwertują sygnały promieniowania optycznego – za pośrednictwem zoptymalizowanych fotodiod – na wynik cyfrowy i realizują ciągły lub wyzwalany pomiar. Czułość naświetlenia można regulować za pomocą wzmocnienia, czasu konwersji i częstotliwości zegara wewnętrznego. Układ osiąga rozdzielczość do 24 bitów przy natężeniu oświetlenia do 2,38 nW/cm² i czasie integracji 64 ms.

Fotodioda GUYA-S12SD

Fotodioda GUYA-S12SD firmy Genicom (Korea) jest tania i łatwo dostępną diodą Schottky'ego czułą w pasmie nadfioletu [1]. Dioda ma obudowę SMD (2,8×3,5×1,8 mm) ze strukturą fotoczujną o powierzchni 0,076 mm². Podstawowe parametry [1]:

- Napięcie wsteczne: 5 V max.
- Zakres widmowy: 240...370 nm
- Czułość maksymalna (352 nm): 0,14 A/W
- Prąd ciemny: 1 nA
- Prąd (dla indeksu UV1): 21 nA typ (prąd offsetu 83 nA)



Poprzednie odcinki znajdują się pod adresem:
<https://ulubionykiosk.pl/media>

- Kąt widzenia: 100°

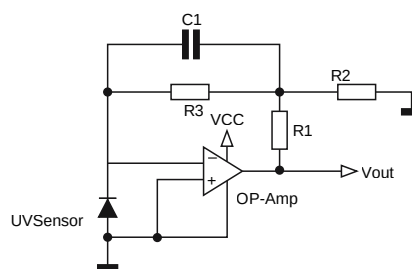
Prąd fotodiody I_{ph} w zależności od indeksu UV można obliczyć ze wzoru:

$$I_{ph} = 21 \cdot UVx + 83[nA]$$

gdzie $x=1...11$

Dla poziomu oświetlenia poniżej UV1 wartość prądu spada liniowo (praktycznie) do zera.

W firmowej nocie aplikacyjnej pokazany jest przykładowy schemat (rysunek 1) układu pomiarowego z fotodiodą GUYA-S12SD [2]. Jest to układ wzmacniacza transimpedancyjnego (TIA) rozbudowany o dzielnik R1 i R2 do postaci sieci T. Taka



Rysunek 1. Przykładowy schemat układu pomiarowego [2]



Fotografia 1. Moduł UV Sensor firmy Waveshare [5]

modyfikacja zwiększa wzmocnienie układu w przybliżeniu o wartość R_1/R_2 , bez konieczności zwiększania wartości rezystora R_3 . Napięcie wyjściowe wynosi w takim przypadku:

$$V_{out} = I \cdot R_3 \cdot \left(1 + \frac{R_1}{R_3} + \frac{R_1}{R_2}\right)$$

Dla $R_1=3,3 \text{ k}\Omega$, $R_2=1 \text{ k}\Omega$, $R_3=1 \text{ M}\Omega$ daje to wzmocnienie prądowe ok. $4,3 \cdot 10^6 \text{ [V/A]}$. Dla oświetlenia UVI1 prąd fotodiody wynosi 104 nA i napięcie wyjściowe $V_{out}=447,2 \text{ mV}$. Zwiększenie oświetlenia do kolejnej wartości indeksu UVI powoduje liniowy wzrost napięcia V_{out} o $90,3 \text{ mV}$.

Moduł UV Sensor

Firma Waveshare udostępnia na swojej stronie Wiki dla modułu UV Sensor instrukcję, schemat i oprogramowanie. Moduł został już wcześniej opisany w artykule „Czujnik światła ultrafioletowego” [3]. Tutaj jednak dokładniej przeanalizujemy jego budowę (fotografia 1). Stopień wstępny jest zrealizowany zgodnie ze schematem z rysunku 1. Zastosowano tu kondensator $C_1=100 \text{ nF}$, co oznacza znaczne ograniczenie pasma częstotliwości oraz szumów. Użyty wzmacniacz SGM8521 firmy SG Micro może pracować przy pojedynczym zasilaniu od $2,1 \text{ V}$ do $5,5 \text{ V}$ z prądem spoczynkowym $5,5 \mu\text{A}$. Zapewnia wejście typu rail-to-rail z szerokim zakresem wejściowego napięcia wspólnego od zera do V_s , bardzo niski wejściowy prąd polaryzacji wynoszący 3 pA i napięcie offsetu $3,5 \text{ mV}$ (maks.).

W module zastosowano też drugi stopień w standardowym układzie wzmacniacza nieodwracającego. Jest to bardzo dobre rozwiązanie przy dołączaniu modułu kabelkami, ponieważ wzmacniacz TIA może zachowywać się niestabilnie przy zmianach pojemności obciążenia. Został tutaj zastosowany jeden wzmacniacz podwójnego układu LM358 firmy Texas Instruments. Układ pracuje z napięciem zasilania od 3 V , wyjściowym napięciem wspólnym od zera do $(V_s-1,5 \text{ V})$ i zakresie napięcia wyjściowego od 5 mV do $(V_s-1,5 \text{ V})$, przy obciążeniu $10 \text{ k}\Omega$. Użyty potencjometr $20 \text{ k}\Omega$ z rezystorem $10 \text{ k}\Omega$ ustawia maksymalne wzmocnienie tego stopnia na ok. 2 V/V . Dla badanego egzemplarza modułu „UV Sensor” wzmocnienie wynosiło $2,108 \text{ V/V}$. Dla oświetlenia UVI1 daje to napięcie wyjściowe $V_{out}=942 \text{ mV}$.

Zaleca się zasilanie modułu napięciem $3,3 \dots 5 \text{ V}$. Na płytce zamontowana jest czerwona dioda LED sygnalizująca zasilanie. Światło tej diody nie zakłóca pracy czujnika UV.

Dołączanie modułu UV Sensor z płytą Enviro Weather

Dołączenie wyjściowego napięcia modułu UV Sensor do wejścia przetwornika ADC modułu Raspberry Pi Pico W (na płytce Enviro Weather firmy Pimoroni [6]) okazało się problemem. Wbudowany przetwornik ADC procesora RP2040 jest tylko 12-bitowy. W celu utrzymania kompatybilności z arytmetyką 16-bitową odczytane słowa danych są przesuwane w lewo o 4 bity. Powoduje to duże skoki wartości dla niewielkich zmian poziomu wejściowego. Brak dokładnego napięcia odniesienia przetwornika powodował duże, skokowe, przypadkowe zmiany odczytu w wyniku impulsowych zakłóceń

przetwornicy zasilającej procesor. Tych błędnych danych nie udało się usunąć nawet filtrem medianowym – stąd konieczność zastosowania zewnętrznego przetwornika ADC.

Przetwornik ADC typu ADS1115

Przetwornik ADS1115 firmy Texas Instruments jest niezwykle popularnym układem scalonym – o bardzo dużych możliwościach, za rozsądną cenę. Ten 16-bitowy konwerter z interfejsem I²C może pracować z zasilaniem $2,0 \dots 5,5 \text{ V}$. Przy bardzo niskim poborze prądu $150 \mu\text{A}$ (tryb ciągłej konwersji) może próbować od 8 Sps do 860 Sps . Ma wbudowane wewnętrzne źródło napięcia odniesienia o niskim dryfcie oraz cztery multipleksowane wejścia unipolarne (lub dwa bipolarne). Unikalną jego cechą jest wyposażenie w wewnętrzny wzmacniacz PGA o bardzo wysokiej impedancji ($>3 \text{ M}\Omega$), co umożliwia ustawienie sześciu zakresów napięcia wejściowego od $\pm 6,144 \text{ V}$ do $\pm 0,256 \text{ V}$. Niewykorzystane wejścia należy dołączyć do zasilania VDD lub masy.

Obsługa szyny I²C jest dość prosta. Układ udostępnia tylko trzy rejestry 16-bitowe oraz 7-bitowy rejestr adresowy. Pod adresem $0x00$ znajduje się rejestr danych, $0x01$ to rejestr konfiguracji, a pozostałe konfiguruje komparator. Każdy transfer rozpoczyna się od wpisania adresu rejestru do rejestru adresowego. Jeśli po nim zostaną przesłane dwa bajty danych, to układ wpisze je do wybranego wcześniej rejestru. Jeśli natomiast po wpisaniu adresu wystąpi stan START (lub RESTART) z ustawionym bitem odczytu, wówczas zostaną odczytane dwa bajty z rejestru o ostatnio ustawionym adresie.

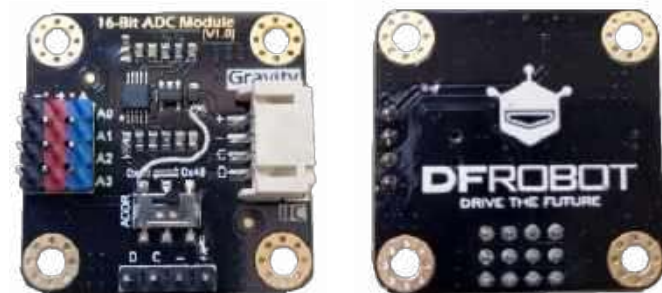
Układ nie umożliwia zgłaszania przerw. Dlatego po wystartowaniu konwersji należy aktywnie odczytywać stan rejestru konfiguracji w oczekiwaniu na zasygnalizowanie zakończenia przetwarzania.

Do pracy z układem ADS1115 zastosowano bibliotekę języka MicroPython [9], opracowaną przez Wolfganga (Wolle) Ewalda w języku Python dla ESP32. Pracuje ona doskonale na RP2040.

Inicjalizacja układu ADS1115 wymaga jedynie wpisania dwóch bajtów do rejestru konfiguracyjnego, co ustawia multiplekser wejść, wzmocnienie PGA, szybkość przetwarzania, tryb pracy (pojedyncza/ciągła) oraz parametry pracy komparatora. Wpis jedynki do bitu OS rozpoczyna przetwarzanie. Podczas odczytu tego rejestru stan bitu OS informuje, czy przetwarzanie jest wykonywane, czy zakończone. Jeśli zostało zakończone – można odczytać dwa bajty z rejestru danych.

Modyfikacja modułu DFRobot I²C ADS1115 do pracy z płytą Raspberry Pi Pico W

Moduł DFRobot I²C ADS1115 (DFR0553) firmy DFRobot [8] zawiera układ ADS1115. Zasilanie układu ADS1115 jest realizowane przez układ LDO typu LP5907MFX-3.3 dostarczający napięcie $3,3 \text{ V}/250 \text{ mA}$. Przy spadku napięcia na tym układzie, wynoszącym od 50 do 250 mV przy obciążeniu – odpowiednio – od 100 do 250 mA , układ można zasiląć napięciem ok. $3,55 \dots 5,5 \text{ V}$. Przedział ten doskonale pasuje do typowego zakresu zasilania z jednego ogniwa Li-Poly.

Fotografia 2. Zmodyfikowana płytka drukowana modułu I²C ADS1115: a) przód, b) tył

Powstaje jednak problem z szyną I²C. Obie linie mają rezystory podciągające 10 kΩ dołączone do napięcia wejściowego. A szyna I²C modułu EnviroWetaher/Grow ma rezystory podciągające 10 kΩ dołączone do własnego zasilania 3,3 V procesora. Aby uniknąć konfliktu i nadmiernego obciążenia prądowego szyny, należy wymontować rezystory podciągające z płytki I²C ADS1115 (**fotografia 2**). Po zdjęciu rezystory są lutowane jednym końcem do ich oryginalnego pola lutowicznego. Ułatwia to późniejsze odtwarzanie pierwotnej funkcjonalności.

Zasilanie płytki I²C ADS1115 jest pobierane z szyny VSYS płytki Raspberry Pi Pico W modułu EnviroWetaher/Grow. Szyna jest zasilana z akumulatora Li-Poly, jednak występują zakłócenia wsteczne od dołączonego do niej wejścia przetwornicy DC/DC typu RT6154A. Można zdecydowanie poprawić jakość zasilania przetwornika ADS1115 poprzez dodanie na wejściu płytki kondensatora tantalowego o niskiej impedancji wewnętrznej (ESR). Rekomendujemy dokonanie kilku modyfikacji płytki:

- Na dolnej stronie płytki drukowanej należy przeciąć ścieżkę idącą do pinu 4 („+”) gniazdka P1 (fotografia 2b).
- Na górnej stronie płytki należy usunąć rezystory R1 i R2 (fotografia 2a).
- Przewodem należy połączyć nóżkę 1 przełącznika S1 (oznaczenie 0x49) oraz nóżkę katody diody D1 (fotografia 2a).
- Najlepiej dodać kondensator 4,7 μF (tantal) wlutowany na nóżkach 3 i 4 gniazdka P2.

Wtedy pin 4 („+”) gniazdka P1 oraz piny 1...4 („+”) gniazdka P4 wyprowadzają regulowane i „czyste” napięcie 3,3 V. Podciąganie linii SCL i SDA należy zrealizować zewnętrznie.

Badania zmodyfikowanej płytki ADS zostały przeprowadzone po podłączeniu do płytki EnviroWetaher. W roli źródła zasilania zestawu użyto akumulatora Li-Ion 18650 o napięciu 4,1 V, podłączonego przewodem o długości ok. 10 cm. Z kolei zasilanie płytki ADS było podłączone do płytki EnviroWetaher przewodem ok. 18-centymetrowym. Natomiast sygnały do wejść oscyloskopu podłączone zostały dosyć długimi przewodami ekranowanymi (bez sondy). Badania zostały wykonane przy relatywnie wysokim poziomie zakłóceń zewnętrznych.

W trakcie wykonywania przez płytkę EnviroWetaher pomiarów z transmisją Wi-Fi występują dwa rodzaje zakłóceń: szpilki oraz impulsy. Szpilki są dosyć dobrze tłumione przez LDO płytki ADS, z poziomu 4,28 mV RMS na szynie VSYS – do 174 μV RMS (2,2 mV pp) na wyjściu LDO. Trochę gorzej jest z impulsami, ale one również są tłumione ponad 10-krotnie.

Występowanie na szynie VSYS płytki Raspberry Pi Pico W (czyli na wejściu przetwornicy DC/DC) trzykrotnie większych zakłóceń impulsowych (o długości ok. 30 μs) niż na jej wyjściu – oraz jednokierunkowy kształt tych zakłóceń – świadczy o problemach ze spadkiem napięcia przy impulsowym poborze prądu (ok. 300 mA).

Pico Inky Pack – moduł z wyświetlaczem e-Paper

Pico Inky Pack (PIM634) firmy Pimoroni to moduł z czarno-białym wyświetlaczem e-Paper o przekątnej 2,9" i rozdzielczości 296×128 pikseli przeznaczony dla Raspberry Pi Pico oraz Raspberry Pi Pico W [10]. Ma wbudowany kontroler, który realizuje komunikację za pomocą interfejsu SPI. Wyświetlacze e-Paper cechują się wysokim kontrastem wyświetlanego obrazu i pobierają prąd

praktycznie tylko w momencie zmiany wyświetlanej treści. Moduł ma wbudowane trzy przyciski. Jest zasilany z szyny 3,3 V modułu Pico.

Pico Graphics to zunifikowana biblioteka do wyświetlania grafiki, dostarczona przez firmę Pimoroni i umożliwiająca sterowanie wyświetlaczami za pomocą Raspberry Pi Pico w języku MicroPython. Udostępnione są również odpowiednie przykłady ułatwiające tworzenie własnych projektów.

Praca modułu UV Sensor z modułem DFRobot I²C ADS1115 i płytką Enviro Weather

Do płytki Raspberry Pi Pico W, przylutowanej techniką powierzchniową do płytki Enviro Weather [6], można wlutować złączę goldpin. Są one usytuowane niejako „na górze” płytki. Umożliwia to dodatkowo pełny dostęp do wszystkich zasobów modułu Raspberry Pi Pico W – w szczególności użycie płytek rozszerzeń. Aby zrealizować ten cel, należy użyć ekspandera, np. Pico Decker (Quad Expander, PIM555) firmy Pimoroni [11]. Jest to ekspander pinów, wyposażony w standardowe złącze żeńskie do bezpośredniego wpięcia modułu RPi Pico lub połączenia za pomocą przewodów. Ma cztery zestawy męskich listew 2×20 pinów, które umożliwiają podłączenie dodatkowych modułów rozszerzeń. Etykiety pinów umieszczone na górnej stronie płytki znacznie ułatwiają prototypowanie.

Aby zmodyfikowaną płytkę Enviro Weather dołączyć do męskich listew ekspandera, trzeba wykonać adapter żeńsko-żeński. Moduł Pico Inky Pack został dołączony bezpośrednio do męskich listew. Natomiast moduły z innymi czujnikami wpięto do ekspandera przy użyciu standardowych kabelków IDC (**fotografia tytułowa**).

Moduł UV Sensor [5] został dołączony do zmodyfikowanego modułu DFRobot I²C ADS1115 (DFR0553) firmy DFRobot [8]. Z płytki ADS zostało dla niego pobrane czyste napięcie zasilania 3,3 V. Zmodyfikowana płytki ADS została dołączona do płytki Enviro Weather firmy Pimoroni [6].

To czyste napięcie jest też dołączone do modułu Grove – I²C Color Sensor opisanego w poprzednim artykule „Czujnik koloru TCS3472 firmy ams-OSRAM” [12] oraz do modułu AS7341 Spectral Color Sensor opisanego w artykule „Czujnik wielospektralny AS7341 firmy ams-OSRAM” [7]. Wszystkie układy scalone zostały dołączone do tej samej szyny I²C.

Zastosowanie płytki Enviro Weather wymaga najpierw wgrania do niej najnowszej wersji oficjalnego pliku obrazu (uf2) – zawierającego MicroPythona oraz biblioteki firmowe, np. Pico Graphics. Następnie należy wpisać folder projektu najnowszej aplikacji Enviro. Procedura ta została dokładnie omówiona w poprzednim artykule „Stacja pogodowa Enviro Weather firmy Pimoroni” [6].

Najprostszy sposób wgrywania oprogramowania firmowego:

1. Pobierz ze strony *Enviro MicroPython firmware* [13] najnowszą wersję firmowego pliku obrazu (uf2) zawierającą jednocześnie Enviro i MicroPythona. W chwili pisania niniejszego artykułu jest to plik *pimoroni-enviro-v1.22.2-micropython-enviro-v0.2.0.uf2*.
2. Trzymaj wciśnięty przycisk BOOTSEL płytki Raspberry Pi Pico W (pod spodem płytki Enviro) i podłącz ją kablem USB do komputera. Spowoduje to przejście oprogramowania płytki

```

ITR=559ALS-01 0x23, PCF85063A 0x51, BME280 0x77, ADS1115 0x48/0x49, AS7341 0x39, TCS34725 0x29
Detected devices at I2C1-addresses:
['0x23', '0x39', '0x39', '0x48', '0x51', '0x77']
[35, 41, 57, 72, 81, 119]
AS7341 initialization
TCS34725 initialization
ADS1115 initialization
2000-01-01 14:46:30 (info / 11kB) - seconds since last reading: 111
Light: 1127
RGB: (72, 80, 94)
UV: 6.187688
2000-01-01 14:46:37 (info / 10kB) > going to sleep for 1 minute(s)
2000-01-01 14:46:37 (debug / 102kB) - clearing and disabling previous alarm
2000-01-01 14:46:37 (info / 100kB) - setting alarm to wake at 14:47pm
2000-01-01 14:46:37 (info / 99kB) - shutting down
2000-01-01 14:46:37 (debug / 97kB) - on usb power (so can't shutdown). Halt and wait for alarm or user reset instead
2000-01-01 14:47:00 (debug / 112kB) - reset

```

Rysunek 2. Informacje w oknie Thonny



Fotografia 3. Pomiar oświetlenia typowej żarówki energetycznej LED

Pico W do trybu DFU i na komputerze otworzy się okno RPI-RP2 pokazujące zawartość dysku udostępnianego przez Pico W.

3. Przeciągnij pobrany plik uf2 do okna dysku RPI-RP2.

Płytkę Pico W uruchomi się ponownie z najnowszą wersją MicroPythona. Nie będzie już udostępniała dysku wymiennego oraz przejdzie bezpośrednio do trybu konfiguracji.

Uwaga! Nastąpi skasowanie poprzedniej zawartości pamięci Flash.

Teraz trzeba do Raspberry Pi Pico W wpisać pliki aplikacyjne z pobranego pliku *enviro_sensors.zip* (<https://tiny.pl/nvhw18p>). Podmianie ulegną pliki *main.py* i *config.py*.

Oprogramowanie uruchamiane było w środowisku Thonny. Odczyt z czujników jest realizowany w pętli, na której końcu płytkę Enviro Weather jest wprowadzana na 60 sekund w uśpienie z aktywnym zewnętrznym układem scalonym zegara RTC. Podczas zasilania z akumulatora wyłączeniu ulega zasilanie całej płytki Enviro Weather oraz wszystkich dołączonych czujników i modułu ADS1115, z wyjątkiem układu RTC. Z tego względu w pętli każdorazowo musi być wówczas wykonywana inicjalizacja układu ADS1115, a potem odczyt danych pomiarowych. Gdy podłączone jest zasilanie z portu USB, procesor RP2040 aktywnie oczekuje na ustawienie alarmu przez RTC – następnie wymuszany jest Reset. Oznacza to utratę łączności środowiska Thonny z płytką Raspberry Pi Pico W. Aby tego uniknąć, można w pliku */Enviro/__init__.py* „zakomentować” linię 646 *machine.reset()*. Działanie takie nie zmienia funkcjonalności aplikacji, natomiast umożliwia ciągły podgląd wydruków w oknie Shell środowiska Thonny, w którym oprogramowanie firmowe Enviro Weather wyświetla dane informacyjne i debugowe (**rysunek 2**). Aplikacja także prezentuje przydatne informacje oraz dane pomiaru UV.

Na początku pętli pomiarowej wykonywana jest detekcja obecności wszystkich układów na szynie I²C. Obsługiwane są tylko aktywne czujniki – dzięki czemu możliwe staje się odłączanie i przyłączanie modułów czujnikowych w trakcie pracy oprogramowania. Zapobiega to również zawieszaniu procesora w przypadku niestabilnego kontaktu przewodów połączeniowych.

Przykład działania całego układu z oświetleniem słonecznym (w cieniu) został pokazany na zdjęciu tytułowym. Zastosowanie dość dużego (i niedrogiego) wyświetlacza e-Paper (który wyświetla dane nawet bez aktywnego zasilania), w połączeniu z dynamicznym wyłączaniem zasilania całości pomiędzy pomiarami, umożliwia uzyskanie minimalnego poboru mocy.

Przykład pomiaru oświetlenia typowej żarówki energetycznej LED zaprezentowano na **fotografii 3**. Na górze wyświetlacza pokazywane są: data i czas. Bez podłączenia Wi-Fi wyświetlana jest data zapisana w pliku *last_time.txt*. Po prawej stronie prezentowane są dane z dziesięciokanałowego czujnika wielospektralnego AS7341 [7], a po lewej stronie – od góry – dane odczytane z czujników płytki Enviro Weather: ciśnienie, temperatura i wilgotność z BME280; poziom oświetlenia z LTR-559. Dalej pokazywane są: FL – detekcja i częstotliwość migotania z AS7341; poziom UV (w mV); CAS – poziom oświetlenia z AS7341, CTCS – poziom oświetlenia

z TCS34725 oraz (w ostatniej linii) kolory RGB z TCS34725. Jeśli nie podano jednostek, to dane są reprezentowane przez surowe wartości odczytu.

Próby pokazały bardzo dokładne i stabilne pomiary napięcia wykonywane przez układ ADS1115.

Podsumowanie

Prezentowane rozwiązanie stanowi próbę szybkiego prototypowania polegającego na rozbudowie systemu Enviro Weather o kolejne czujniki środowiskowe. Dlatego w artykule skupiono się na komunikacji na szynie I²C i obsłudze czujników. Nie skorzystano z możliwości łączenia bezprzewodowego modułu Enviro Weather z chmurą. W kodzie oprogramowania pozostawiono różne opcjonalne rozwiązania. Opis czujników został zamieszczony w poprzednich artykułach [6, 7]. W następnej kolejności do projektu można dołączyć inne wcześniej opisane czujniki, np. CO, CO₂, deszczu i pyłów.

Henryk A. Kowalski
Instytut Informatyki
Politechnika Warszawska

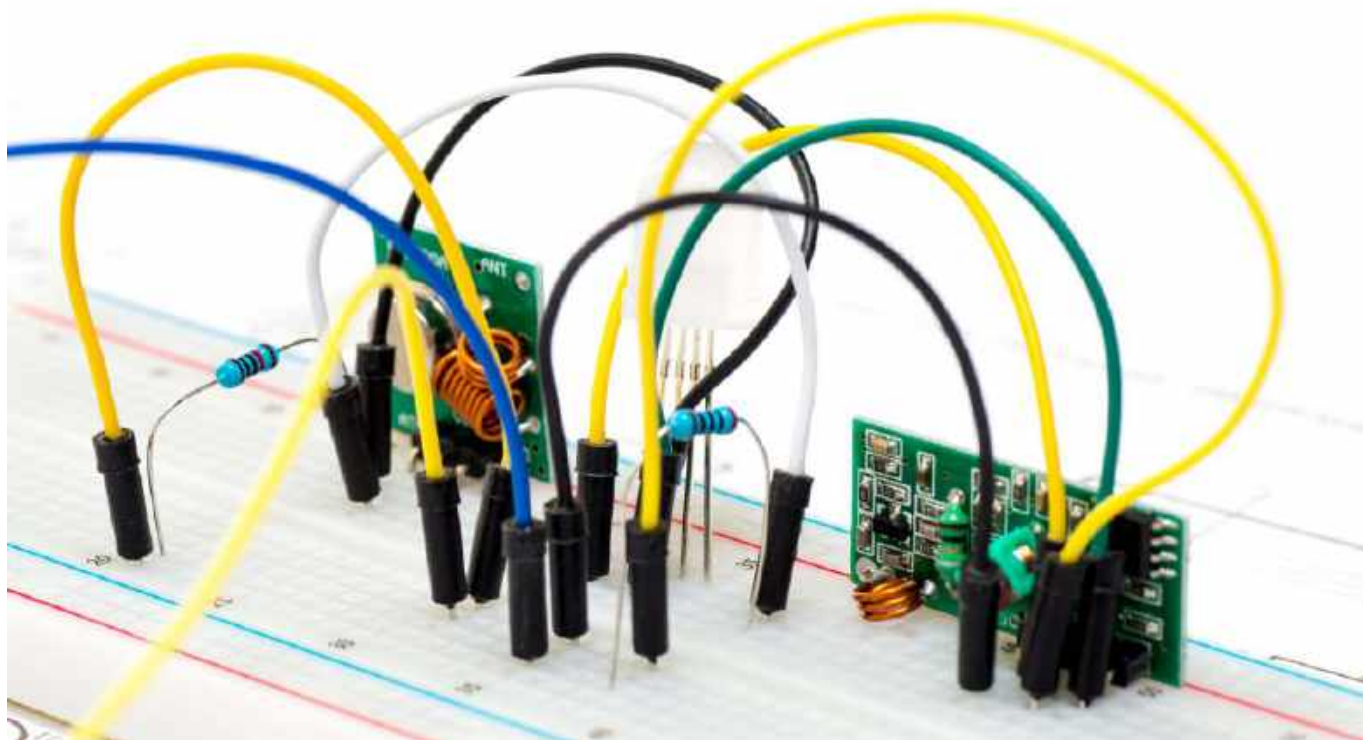
Literatura

- [1] GUV-A-S12SD, Data Sheet, Genicom, <https://tiny.pl/h3-q0q5n>
- [2] GUV-A-S12SD, Application Note, Genicom, <https://tiny.pl/633grpwp>
- [3] Czujnik światła ultrafioletowego, Ryszard Szymaniak, EP 12/2020, <https://tiny.pl/2mm1j3gq>
- [4] Czujniki optyczne (4). Sensory koloru oraz czujniki i moduły multispektralne, Przemysław Musz, EP 10/2023, <https://tiny.pl/dnpjk>
- [5] UV Sensor, SKU: 9537, Waveshare, <https://tiny.pl/84f8j5tb>
- [6] Stacja pogodowa Enviro Weather firmy Pimoroni, Henryk A. Kowalski, EP 4/2024, <https://tiny.pl/d93r1>
- [7] Czujnik wielospektralny AS7341 firmy ams-OSRAM, EP 8/2024, <https://tiny.pl/ny88mxyx>
- [8] Gravity: I²C ADS1115 16-Bit ADC Module, DFR0553, DFRobot, <https://tiny.pl/n6qfj60b>
- [9] ADS1115_mpy, A MicroPython module for the ADS1115 ADC. Wolfgang (Wolle) Ewald, <https://tiny.pl/twv9zkk9>
- [10] Pico Inky Pack PIM634, Pimoroni, <https://tiny.pl/94b0pv9f>
- [11] Pico Decker (Quad Expander) PIM555, Pimoroni, <https://tiny.pl/mwp5n68n>
- [12] Czujnik koloru TCS3472 firmy ams-OSRAM, EP 9/2024, <https://tiny.pl/7y4vjwms>
- [13] Enviro MicroPython firmware, Pimoroni, <https://tiny.pl/dt49f>

REKLAMA

Mnóstwo doskonałych projektów, tylko na:

EP.com.pl



Prototypowe sztuczki i kruczki

Prototypy są po to, by popełnić na nich tzw. błędy młodości. Jednak pewnej części pomyłek z powodzeniem da się uniknąć jeszcze przed postawieniem pierwszej kreski na ekranie komputera. O czym mowa i w jakie pułapki lepiej nie dać się złapać? A jeśli już dojdzie do błędu – to jak sobie z nim poradzić?

Układy prototypowe służą do tego, by testować na nich pewną ideę, koncepcję, sposób rozwiązania jakiegoś problemu. Inwestor chce, by prototyp zadziałał jak najszybciej. Księgowość każe się liczyć z budżetem. Z kolei projektant chce poświęcić konstrukcji możliwie mało czasu, by nie poprawiać jej wielokrotnie i móc przejść do kolejnego zlecenia. Wszystkie wymienione oczekiwania można spiąć jedną klamrą: potrzebny jest dobry projekt. Im więcej kwestii uwzględnimy na etapie projektowania, czym lepiej dopracujemy nasze urządzenie jeszcze na etapie wyobraźni i obliczeń, tym łatwiej będzie sprostać wszystkim trzem wymaganiom jednocześnie.

W moim katalogu błędów projektowych wyróżniam dwa rodzaje: da się go poprawić drutem – albo nie da się go poprawić drutem. Jeżeli inwestor nie będzie widział jakiegoś „druciaka” w czasie testów prototypu, a płynące prądy są małe lub zbiega wolne, to znaczy, że określoną pomyłkę da się skorygować drutem. Zwłaszcza jeżeli całość usztywni się później poprzez zalanie fragmentu płytki klejem epoksydowym – tak, aby żadne połączenie nie oderwało się od PCB. Mój wykładowca od układów elektronicznych mawiał nawet, że prototyp bez drutu jest jakiś taki... wybrakowany.

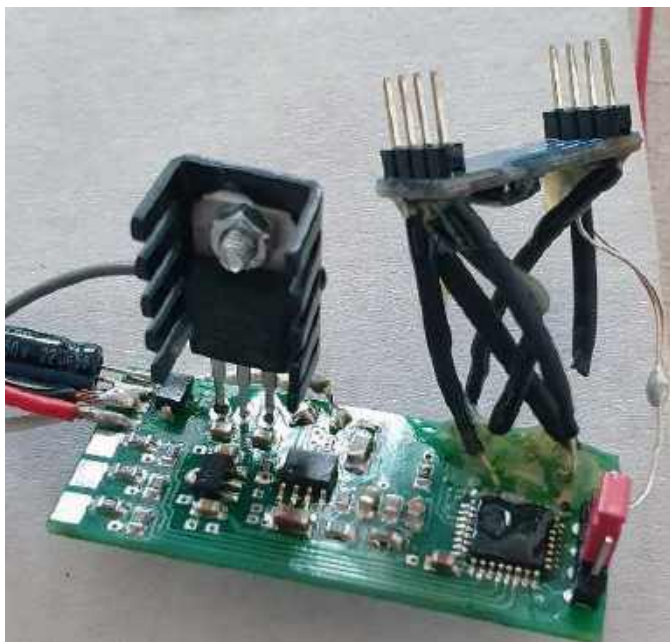
Mój osobisty faworyt w tym zakresie to cienka srebrzanka (mam kilka szpilek o różnych średnicach), ponieważ łatwo można ją dołutować nawet wprost do wyprowadzenia układu scalonego czy innego podzespołu. Do obudów z wąskimi nóżkami, jak TQFP64,

który jest już przylutowany do powierzchni płytki, stosuję z kolei pojedyncze miedziane włoski, uzyskane w wyniku zdjęcia izolacji z jednej żyły typu linka. Taki włoszek wystarczy tylko pocynować, po czym przyłożyć do odpowiedniego wyprowadzenia i docisnąć gorącym grottem – **fotografia 1**. Można w ten sposób uratować prototyp, kiedy jakieś wyprowadzenie układu scalonego zostało pomyłone – lub jego funkcja źle zinterpretowana. Miedziane włoski mają tę przewagę nad srebrzanką, że są od niej elastyczniejsze, co utrudnia odłamanie ich od miejsca przylutowania, które ma przecież niewielką powierzchnię. Na koniec warto owo miejsce wyczyścić (choć zgrubnie) z kalafonii i zalać żywicą epoksydową w celu wzmocnienia.

Jeżeli jednak w grę wchodzi bardzo dokładne kontrolowanie impedancji, zwłaszcza linii różnicowej, to lepiej już nie podchodzić do takiego błędu z drutem – tak samo, jak przy pomyleniu obudowy całego układu scalonego, brrrr... Na **fotografii 2** widać układ prototypowy, w przypadku którego projektant (tak, to ja) zbyt pochopnie podszedł do układu W25Q16JVSNIQ, a konkretniej do jego obudowy. Po wpisaniu wymienionego symbolu w przeglądarkę otrzymujemy strony dystrybutorów, na których widać zdjęcia układu w typowej, wąskiej



Fotografia 1. Układ w obudowie TQFP64 z dodatkowym doprowadzeniem przylutowanym bezpośrednio do nóżki



Fotografia 2. Układ prototypowy po pomyleniu rozmiaru obudowy jednego z układów scalonych

obudowie SO8 150 mils, ba – często jest ona opisana nawet jako SO8. Co zrobił Michał? Michał postawił na płytce ten układ w obudowie SO8, bo nie kwapił się, by przeczytać od deski do deski notę katalogową nieznanego sobie układu. Co potem robił Michał w nocy? Lutował srebrząnką zdobytą naprędce płyteczkę – zawierającą ten układ – do padów na gotowej już płytce, bo termin goni, a w Polsce zdobycie tego układu w obudowie SO8 (150 milsów) jest niemożliwe. Za to egzemplarzy w obudowie SOP8 (200 milsów) jest na półce i takiego właśnie chciał ten nieuważny Michał użyć. Ważne, że żywica epoksydowa utrzymała wspomniane elementy w całości. Przy bardziej zaawansowanym układzie scalonym, wymagającym np. dokładnego odsprężania zasilania, podobna sztuczka okazałaby się już niemożliwa do wykonania.

Konkluzja: jeżeli nie znamy jakiegoś elementu, przeczytajmy jego notę katalogową pod kątem oznaczenia handlowego danej wersji obudowy. Można się srogo zdziwić, bo dystrybutorzy nierazko mieszają zamieszczane zdjęcia poglądowe obudów, a nawet nazwy tychże. Zauważyłem, że szczególnie często dotyczy to układów w obudowach SO (150 milsów) i SOP (200 milsów), a przecież nie są one kompatybilne. W ogóle, z oznaczeniami SO, SOP, SOIC, TSOP, SSOP itd. mamy niezłą „jazdę po bandzie” – dlatego zachęcam do uważnego analizowania rysunków obudów wraz z ich wymiarami, bo każdy producent może stosować nieco inne standardy w zakresie dokumentacji.

Kiedy mamy już dopracowany projekt, warto poświęcić więcej uwagi możliwości jego stopniowego uruchamiania, dokładniej:



Rysunek 1. Rezystor 0 Ω włączony szeregowo z linią zasilającą układ scalony

uruchamiania poszczególnych bloków w obrębie jednej płytki. Jeśli na przykład mamy na płytce zawarte fikuśne zasilacze, blok analogowy, przetwarzanie analogowo-cyfrowe i blok cyfrowy, warto byłoby sprawdzić, czy z zasilaczami jest wszystko w porządku, zanim podłączy się do nich bardzo drogą i delikatną cyfrówkę. Pół biedy, gdy można sobie pozwolić na montowanie płyty fragmentami. Gorzej, jeśli zawiera ona masę podzespołów w małych obudowach o gęstym rastrze wyprowadzeń, co więcej: z wkładkami do odprowadzania ciepła. Polutowanie czegoś takiego ręcznie jest bardzo trudne i czasochłonne, na dodatek można wiele z tych układów po prostu przegrzać gorącym powietrzem. Trzeba więc zastosować montaż całości na maszynie pick & place, a potem lutowanie w piecu.

W takich sytuacjach staram się stosować rozwiązanie jak na widoku płytki z **rysunku 1**. Między wyjściem zasilacza a wejściem zasilania do danego podzespołu (bądź bloku podzespołów) włączam rezystory SMD o niewielkim rozmiarze – tutaj jest to 0805 – lecz nie zlecam ich przylutowania w fabryce. Mając wyjście zasilacza skupione w jednym polu lutowniczym, mogę do niego podłączyć woltomierz, oscyloskop, sztuczne obciążenie lub rezystory obciążające – wszystko, co tylko chcę, by sprawdzić poprawność działania tego modułu.

Kiedy poprzedni blok (w tym przypadku stabilizator) jest już w pełni sprawny i uruchomiony, wlutowuję ręcznie rezystor 0 Ω, będący tak naprawdę zwykłą zworą, co umożliwia zasilanie docelowego już układu. Zajmuje to minutę lub dwie, a moja głowa jest spokojna, że na testowany blok trafiają prawidłowe napięcia. Płytki w wersji produkcyjnej zostały pozbawiane wspomnianych rezystorów-zwor, choć przyznam, że nie zawsze – kiedy zdarzyło mi się zapomnieć o ich usunięciu i powstała niemała partia płytek zawierająca te (niepotrzebne już na tym etapie) dodatki, zleciłem po prostu wlutowanie w ich miejsce rezystorów 0 Ω. Być może ktoś, kto w przyszłości chciałby przeprowadzić reverse engineering, będzie miał potężną zagwozdkę, co autor miał na myśli...

Michał Kurzela, EP

REKLAMA

Mnóstwo doskonałych projektów, tylko na:

EP.com.pl

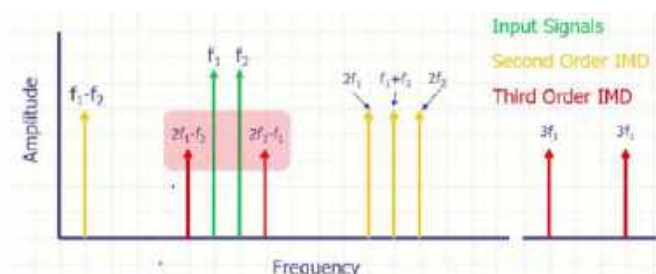
Amatorskie pomiary wzmacniaczy audio (4)

W czwartym odcinku naszego cyklu przyjrzymy się kolejnym, bardzo istotnym zagadnieniom związanym z obiektywnymi metodami testowania wzmacniaczy audio. Tym razem zajmiemy się tematyką zniekształceń intermodulacyjnych oraz dynamicznych.

Pomiary zniekształceń intermodulacyjnych

Zwykły pomiar zniekształceń harmoniczných wytwarzanych przez jednotonowe fale sinusoidalne, nawet o różnych częstotliwościach, nie dostarcza wielu informacji wymaganych do oceny potencjalnej jakości wzmacniacza. Układ powinno się też oceniać pod kątem zniekształceń intermodulacyjnych (IMD) wytwarzanych po podaniu na wejście jednocześnie dwóch lub więcej określonych tonów. Co prawda tego typu zniekształcenia są badane głównie w kanałach komunikacyjnych wielkiej częstotliwości, bo tam stają się bardzo uciążliwe, jednak testowanie wzmacniaczy audio jest również praktykowane i użyteczne.

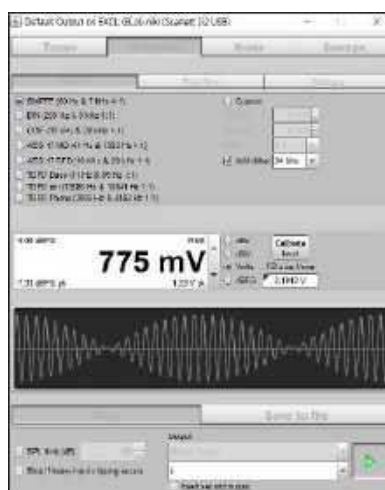
IMD powstaje, gdy dwa lub więcej tonów audio intermoduluje ze sobą w urządzeniu nieliniowym, tworząc niepożądane nowe tony. Podstawowym mechanizmem wytwarzającym IMD w większości urządzeń jest AM (modulacja amplitudy), która tworzy wstępni boczne, stanowiące sumę i różnicę częstotliwości oryginalnych tonów audio. Na przykład, jeśli tony audio o częstotliwościach 2 kHz i 7 kHz przechodzą przez urządzenie nieliniowe i ulegną zniekształceniom intermodulacyjnym AM, sygnał wyjściowy będzie zawierał nowe tony (produkty modulacji) przy 9 kHz (suma tonów) i przy 5 kHz (różnica tonów). Produkty modulacyjne mogą również intermodulować ze sobą i z oryginalnym sygnałem audio, tworząc jeszcze więcej produktów modulacji.



Rysunek 35. Powstawanie zniekształceń intermodulacyjnych IMD

Na rysunku 35 pokazano, jak powstają zakłócenia intermodulacyjne z dwóch tonów podstawowych o częstotliwościach f_1 i f_2 . Najbardziej niepożądane są sygnały IMD 3. rzędu o częstotliwościach $2f_1-f_2$ i $2f_2-f_1$, bo znajdują się zawsze w paśmie użytecznym i nie można ich odfiltrować.

Pomiar zniekształceń IMD zaczynamy od ustawienia generatora, tak by generował odpowiedni sygnał dwutonowy.



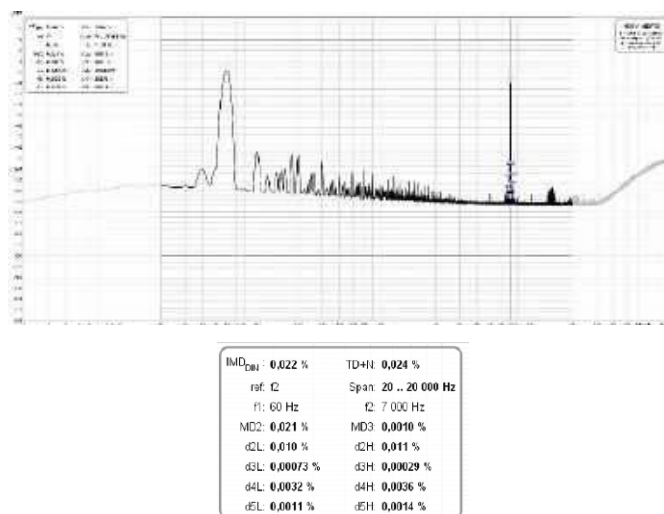
Rysunek 36. Generator sygnału dwutonowego do pomiaru zniekształceń IMD



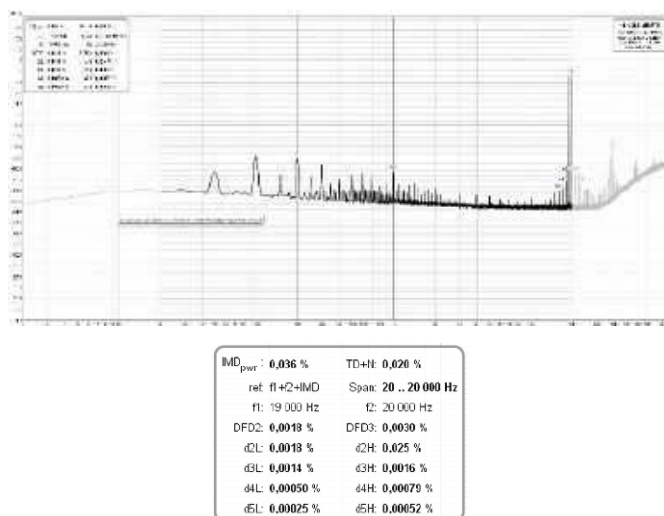
Poprzednie odcinki znajdują się pod adresem:
<https://ulubionykiosk.pl/media>

W tym celu klikamy zakładkę Multitone, a potem Dual Tones – rysunek 36.

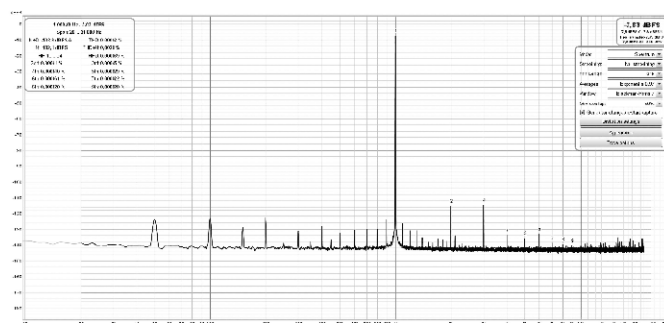
Można tu wybrać kilka standardów częstotliwości oraz stosunku amplitud. Jedną z najbardziej znanych i chyba najstarszą jest metoda SMPTE. Sygnał testowy składa się z tonów o niskiej częstotliwości



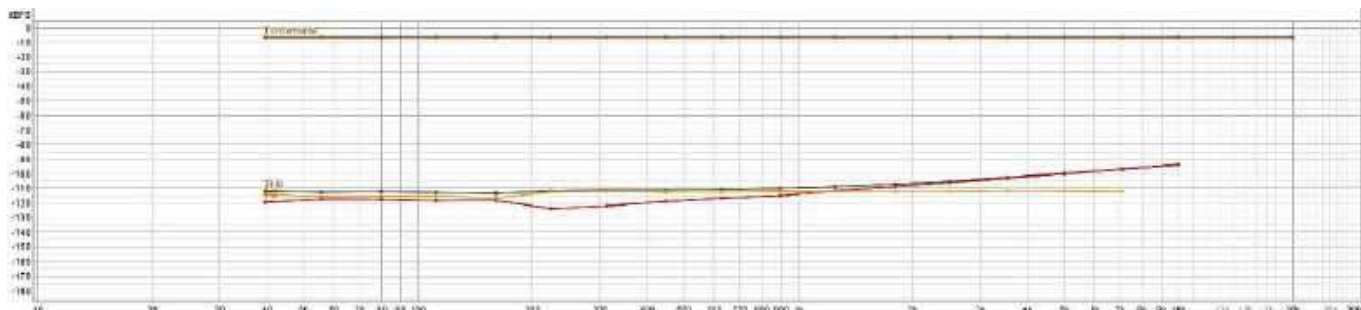
Rysunek 37. Pomiar zniekształceń IMD metodą SMPTE



Rysunek 38. Pomiar IMD metodą CCIF



Rysunek 39. Pomiar DAC zaopłonego z wejściem karty



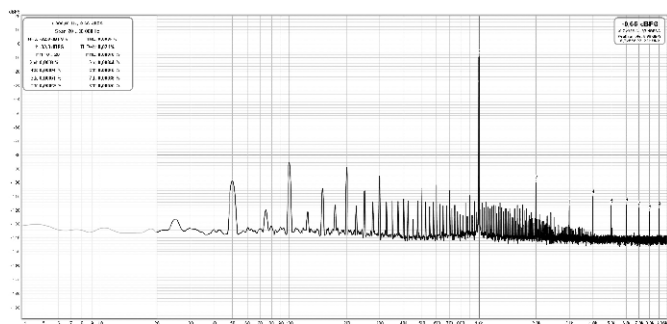
Rysunek 40. Charakterystyka THD w funkcji częstotliwości dla zapętlonej karty z wyjściem DAC

(zwykle 60 Hz) i wysokiej częstotliwości (zwykle 7 kHz), zsumowanych razem w stosunku amplitudy 4 do 1. Inne współczynniki amplitudy i częstotliwości (np. DIN 250 Hz + 8 kHz) są używane sporadycznie. Sygnał SMPTE jest podawany do testowanego urządzenia, a sygnał wyjściowy bada się pod kątem modulacji składowej o górnej częstotliwości przez ton niskiej częstotliwości. Drugi ważny test określany jest skrótem CCIF. Test zniekształceń CCIF IM różni się od testu SMPTE tym, że do badanego urządzenia przykładana jest para tonów o zbliżonej częstotliwości i tych samych amplitudach. Nieliniowość w testowanym urządzeniu powoduje intermodulację między dwoma sygnałami, które są następnie mierzone. Dla typowego przypadku sygnałów wejściowych, przy 19 kHz i 20 kHz, składowe intermodulacyjne będą miały częstotliwości 1 kHz, 2 kHz, 3 kHz, 4 kHz, 5 kHz itd. oraz 18 kHz, 21 kHz, 17 kHz, 22 kHz, 16 kHz, 23 kHz itd.

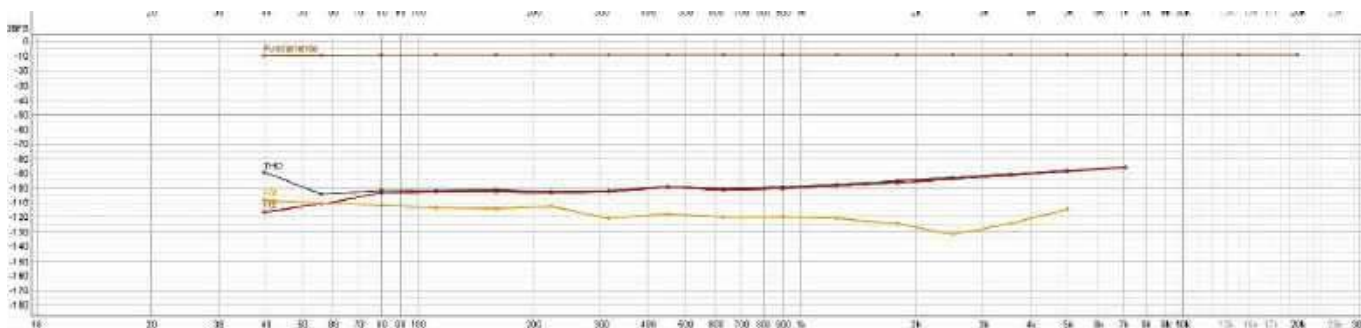
Okno RTA dla pomiaru IMD metodą SMPTE dla mocy wyjściowej 2 W zostało pokazane na **rysunku 37**. Oprócz wartości zniekształcenia IMD wyliczane jest również zniekształcenie Multi-tone Total Distortion Plus Noise, czyli TD+N dla dwóch tonów.

Mamy tu $IMD=0,022\%$ i $TD+N=0,024\%$. To dobry wynik. Przypomnijmy, że $THD+N$ wyniósł $0,0036\%$, ale wzmacniacz przy pomiarze IMD pracuje w trudniejszych warunkach.

Na **rysunku 38** pokazano pomiar IMD metodą CCIF. IMD jest tutaj większe, bo na poziomie $0,036\%$. Na wykresie widać składową d2L o częstotliwości 1 kHz, będącą wynikiem różnicy częstotliwości testowych 20 kHz i 19 kHz. Test CCIF ma tę zaletę, że generuje zniekształcenia w paśmie słyszalnym.



Rysunek 41. Pomiar THD+N wzmacniacza LM1876 przy mocy 2 W



Rysunek 42. Charakterystyka THD w funkcji częstotliwości dla wzmacniacza LM1876 przy mocy 2 W

Zmiana konfiguracji pomiarowej

Zastosowanie wyjścia słuchawkowego karty Scarlett 2i2 tylko częściowo rozwiązało problem pomiarów $THD+N$. Układ dobrze się mierzy, nic się nie wzbudza, ale przy niskich częstotliwościach (ok. 50 Hz) pomiary zniekształceń samej karty okazują się nieco zawyżone. Nie są to duże wartości i można próbować je korygować, jeżeli zależy nam na dokładnych wynikach. Osobiście uważam, że w amatorskich pomiarach nie ma to większego znaczenia. Jak już wspominałem, można by spróbować desymetryzacji wyjściowego sygnału symetrycznego, ale to również może stać się źródłem zniekształceń. Kiedy szukałem możliwości pomiarów parametrów przetworników cyfrowo-analogowych, wpadłem na pomysł rozwiązania tego problemu. Pomocne będą tutaj drivery Java użyte w programie REW. Umożliwiają one niezależne sterowanie urządzeniami wyjściowymi (głośniki) i wejściowymi (mikrofony). Można zatem jako analizator (wejście) zastosować kartę Scarlett 2i2, a jako wyjście generatora (głośnik) – jakiś inne urządzenie. Urządzeniem wyjściowym (generatorem) został przetwornik cyfrowo-analogowy z układem PCM1794A i wejściem USB. Przetwornik mojej konstrukcji ma wyjście asymetryczne i teoretycznie powinien się nadawać do tego celu, bo akceptuje częstotliwości próbkowania do 192 kHz, podobnie jak konwerter USB/I²S. Teoretycznie, bo w praktyce okazało się, że drivery użyte w REW z tym konwerterem USB/I²S da się w tym przypadku skonfigurować prawidłowo tylko dla częstotliwości 44,1 kHz i zależy to od oprogramowania konwertera. Wszelkie inne ustawienia dawały dziwne wyniki w trakcie kalibracji wyjścia przetwornika zapętlonego z wejściem karty Scarlett 2i2. Nie przeszkadza to w podstawowych pomiarach, ale ogranicza możliwości w paśmie powyżej 7...8 kHz – z powodów, które już opisywałem. Trzeba też pamiętać o tym, że zegary generatora (DAC) oraz analizatora (wejście karty) nie są z sobą zsynchronizowane. Taka sytuacja powoduje pewne ograniczenia w konfiguracji pomiarów. Dokładne informacje można znaleźć w plikach pomocy programu REW.

Na **rysunku 39** pokazano zniekształcenia zapętlonego wyjścia DAC z wejściem symetrycznym karty.

Otrzymujemy THD równe $0,00063\%$ i $THD+N$ równe $0,0020\%$. To wyniki lekko gorsze niż przy zapętlaniu wyjść/wejść symetrycznych karty i porównywalne przy zapętlaniu wyjścia słuchawkowego SE karty z wejściem symetrycznym. Przy okazji te wyniki można traktować jako pomiar samego DAC dla 1 kHz. W tym

momencie można sobie zadać pytanie, co zyskujemy, stosując DAC, skoro wyniki dla wyjścia słuchawkowego są tak bardzo podobne? Różnica staje się widoczna przy małych częstotliwościach. Pokazaną na **rysunku 40** charakterystykę pomiaru THD+N w funkcji częstotliwości można porównać z rysunkiem 32.

Teraz zniekształcenia dla małych częstotliwości są wyraźnie niższe i mają równiejszy przebieg. Widać też, że rosną wraz z częstotliwością. Oczywiście pomiar został ograniczony do 7 kHz, przy której to wartości zarejestrowano jeszcze 3. harmoniczną.

Wyniki pomiarów wzmacniacza w układzie pomiarowym z DAC dla mocy 2 W zostały pokazane na **rysunkach 41 i 42**.

Mamy więc tu charakterystykę w funkcji częstotliwości toru pomiarowego wyraźnie równiejszą, a pomiary – bardziej wiarygodne.

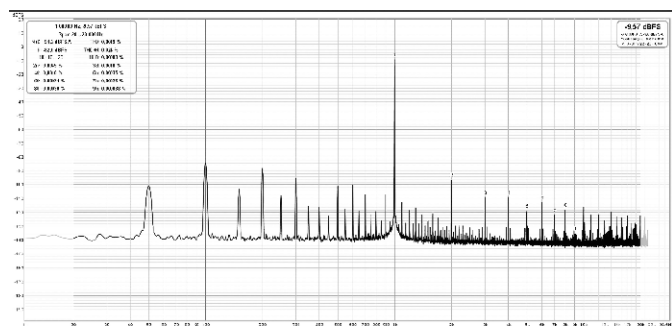
Pomiar maksymalnej mocy wyjściowej

Moc wyjściową mierzy się przy znamionowym obciążeniu rezystancyjnym. Wzmacniacze są najczęściej przystosowane do obciążenia o wartości 8 Ω i/lub 4 Ω . Maksymalną moc wyjściową określa się dla najwyższego dopuszczalnego poziomu zniekształceń nieliniowych THD. Dla dobrych wzmacniaczy mogą to być na przykład zniekształcenia na poziomie 1%. Szybki wzrost zniekształceń nieliniowych w okolicy maksymalnej mocy wyjściowej jest spowodowany obcinaniem szczytów wzmacnianego sygnału z powodu ograniczenia wydajności prądowej układu zasilania.

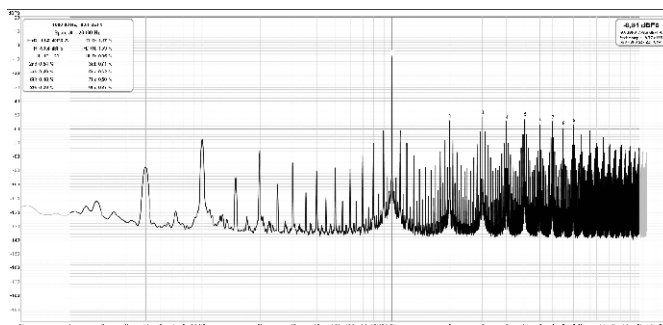
Dokładny pomiar mocy wyjściowej przy założonym poziomie zniekształceń nieliniowych wymaga układu mierzącego te zniekształcenia. Używa się też metody szybkiego, zgrubnego pomiaru. Wzmacniacz obciążamy rezystorem sztucznego obciążenia. Na wejście podajemy regulowany sygnał sinusoidalny z generatora (o częstotliwości 1 kHz), a do wyjścia wzmacniacza podłączamy oscyloskop. Zwiększamy amplitudę na wejściu do momentu, kiedy na ekranie oscyloskopu zaczynamy obserwować początek obcinania szczytów sinusoidy. Napięcie V_{rms} – zmierzone tuż przed początkiem obcinania szczytów – określa nam, przy zadanej rezystancji obciążenia, jaka jest moc wyjściowa. Jak już wspominałem, ta metoda jest mało dokładna. Kiedy już zaczynamy obserwować zmiany kształtu, zniekształcenia są na poziomie 1...2%. Żeby nieco ten pomiar uwiarygodnić, zmniejsza się „trochę” napięcie wyjściowe i przyjmuje, że jest ono właściwe dla mocy maksymalnej wzmacniacza.

Zmierzymy teraz moc naszego wzmacniacza przy użyciu oscyloskopu i układu mierzącego zniekształcenia. Najpierw podłączamy wyjście wzmacniacza do układu sztucznego obciążenia, a sygnał na wejście wzmacniacza jest podawany z wyjścia karty. Potem, za pomocą potencjometru, ustawiamy poziom sygnału na wejściu wzmacniacza – tak by napięcie na jego wyjściu miało wartość trochę mniejszą od zaobserwowanego obcinania szczytów sygnału wyjściowego. W naszym przypadku było to ok. 11 V (RMS). Moc wyjściowa wzmacniacza wynosiła wtedy ok 15 W. Charakterystyka zniekształceń w tym przypadku wygląda tak, jak na **rysunku 43**.

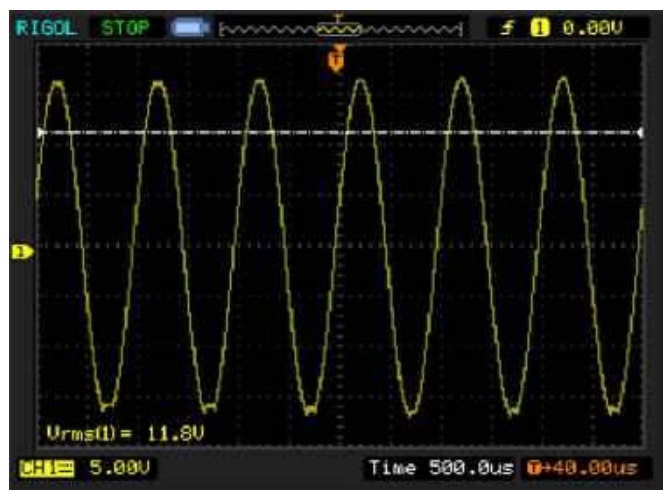
Zniekształcenia utrzymują się nadal na niskim poziomie: THD jest równe 0,0049%, a THD+N wynosi 0,024%. Teraz zwiększamy



Rysunek 43. Zniekształcenia wzmacniacza dla mocy ok. 15 W



Rysunek 44. Zniekształcenia ok. 1,5% dla mocy maksymalnej wzmacniacza



Rysunek 45. Przebieg wyjściowy przy napięciu skutecznym 11,8 V

napięcie na wyjściu wzmacniacza i obserwujemy zniekształcenia. Jeżeli będą miały zadany poziom, to mierzymy napięcie na wyjściu i na jego podstawie wyliczamy moc maksymalną. Ja założyłem, że będzie to 1,5% – **rysunek 44**.

Napięcie wyjściowe wyglądało tak, jak na **rysunku 45**. Widać niewielkie, ale zauważalne obcinanie dolnej połówki przebiegu. Napięcie VRMS wynosi 11,8 V, czyli maksymalna moc wyjściowa wzmacniacza dla THD+N=1,5% to ok. 17,4 W.

Metoda obserwacji przebiegu wyjściowego również da wynik, który może być akceptowalny, ale nie dokładny. Mój oscyloskop ma rozdzielczość pionową 8 bitów i obserwowanie obcinania przebiegów jest utrudnione. Lepszy do tego celu byłby dobry oscyloskop analogowy lub nowszy oscyloskop cyfrowy o pionowej rozdzielczości np. 12 bitów.

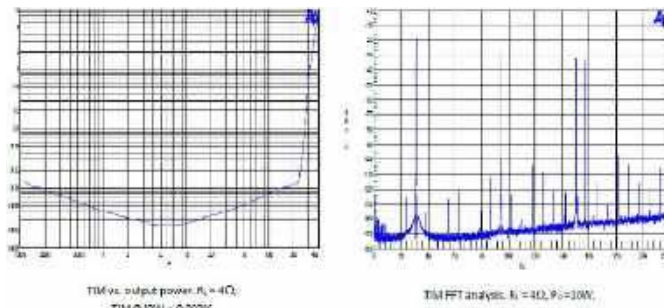
Pomiary zniekształceń dynamicznych

Pomimo że pomiary zniekształceń THD i THD+N są przydatne w badaniu jakości wzmacniaczy, to niestety mają również ograniczenia. Pierwsze z nich wiąże się z faktem, że są to pomiary statyczne. Wzmacniacz pracuje z sygnałem wejściowym o stałej amplitudzie i częstotliwości – i jest to praca bardzo komfortowa w stanie ustalonym. Nie byłoby trudno zbudować wzmacniacza, który będzie miał niskie lub bardzo niskie THD, ale jego brzmienie okaże się takie sobie, żeby nie powiedzieć: fatalne. Przykładem są wzmacniacze tranzystorowe z początku lat 70. XX wieku. Stosowano w nich wyjściowe układy quasi-komplementarne z dwoma tranzystorami mocy o polaryzacji NPN. Dodatkowo tranzystory mocy miały małą częstotliwość graniczną: na przykład popularny 2N3055 ma częstotliwość $f_T=2$ MHz (2 MHz przy wzmocnieniu 1). Żeby taki stopień miał akceptowalne zniekształcenia nieliniowe THD, trzeba było go obciążyć silną pętlą ujemnego sprzężenia zwrotnego. To – w połączeniu z niezbyt szybkimi tranzystorami mocy – stanowiło dobry przepis na powstawanie przykrych, wyraźnie słyszalnych zniekształceń.

I to mimo faktu, że THD kształtowało się na poziomie dzisiejszych dobrych wzmacniaczy. Dlatego wielu uważa, że przy wyborze wzmacniacza nie powinno się kierować wyłącznie parametrem THD+N, często podawanym tylko dla 1 kHz. Lepszym, bo trudniejszym do spełnienia kryterium, są niskie zniekształcenia IMD.

Ale żeby wzmacniacz dobrze brzmiał, musi też mieć odpowiednie parametry dynamiczne. Należy zmusić go w czasie pomiaru do takiej pracy, w której będzie musiał dynamicznie „doregulowywać” się do poziomu sygnału z użyciem układu pętli ujemnego sprzężenia zwrotnego. Okazuje się, że niepoprawnie skonstruowane wzmacniacze w dynamicznych stanach generują bardzo przykre dla naszego słuchu zniekształcenia zwane TIM. Nawet niewielka ich wartość – na poziomie np. 0,2% – jest słyszalna, podczas gdy THD na poziomie 0,2% jest praktycznie niezauważalne. TIM, czyli Transient Intermodulation Distortion (przejściowe zniekształcenia intermodulacyjne), występują we wzmacniaczach tranzystorowych korzystających z ujemnego sprzężenia zwrotnego. Opóźnienia sygnału w torze wzmacniacza sprawiają, że – w pętli sprzężenia – nie jest on w stanie prawidłowo skorygować zniekształceń pod wpływem szybkich, przejściowych sygnałów. Pierwsze wzmacniacze tranzystorowe były „wolne”, czyli reagowały zbyt wolno (z opóźnieniem) na szybko narastające sygnały wejściowe. Przy tym – jak już powiedzieliśmy – żeby zapewnić niskie zniekształcenia harmoniczne, były obejmowane silnym ujemnym sprzężeniem zwrotnym. Wspomniane już opóźnienia powodowały, że (przy szybko narastającym napięciu) sygnał z wyjścia nie trafiał w odpowiednim momencie na wejście przez pętlę sprzężenia zwrotnego i cały wzmacniacz w tych momentach pracował bez odpowiedniego sygnału sprzężenia, w związku z czym był chwilowo mocno przesterowany.

Od czasu wykrycia zniekształceń TIM i źródła ich pochodzenia przez Matti Otała, znane są sposoby minimalizowania zjawiska powstawania TIM do akceptowalnego poziomu. Wszystkie stopnie wzmacniające powinny mieć równą charakterystykę, tak by pętla ujemnego sprzężenia zwrotnego nie musiała być zbyt silna. Wzmacniacz z kolei musi być szybki, czyli powinien bardzo szybko reagować swoim wyjściem na strome zbocza narastające sygnałów wejściowych. Dzięki temu zminimalizowane zostaną opóźnienia sygnału. Parametrem określającym szybkość wzmacniacza jest Slew Rate, wyrażany w woltach na mikrosekundę.



Rysunek 46. Pomiary zniekształceń TIM dla wzmacniacza ICEPower 250ASX250

Na przykład wartość 25 V/μs oznacza, że – po przyłożeniu na wejście impulsu o bardzo krótkim czasie narastania – na wyjściu wzmacniacza osiągniemy 25 V po czasie 1 μs. Szybkość narastania sygnału jest związana z pasmem przenoszenia. Jeden z warunków eliminacji zniekształceń TIM stanowi szerokie pasmo przenoszenia, wynoszące przynajmniej 100 kHz dla sygnałów akustycznych w paśmie 20 Hz...20 kHz. Można przyjąć, że nowoczesne wzmacniacze są pozbawione problemu zniekształceń TIM. Dość łatwo da się wykonać wzmacniacz z odpowiednio wysokim slew rate, dużym zapasem pasma przenoszenia, a jednocześnie z niskim globalnym sprzężeniem zwrotnym.

Oczywiście istnieją metody pomiaru TIM. Na wejście wzmacniacza podaje się sygnał sinusoidalny modulowany sygnałem prostokątnym. Według standardu IEC sygnał sinusoidalny ma częstotliwość 15 kHz, a prostokątny 3,15 kHz. My w naszym układzie pomiarowym nie mamy możliwości zmierzenia zniekształceń TIM, ale możemy zmierzyć pasmo wzmacniacza metodą opisaną w następnym odcinku naszego cyklu. Możemy też próbować, choć nie jest to łatwe, określić z grubsza slew rate za pomocą generatora sygnału prostokątnego i oscyloskopu. Mamy też możliwość zmierzenia zniekształceń IMD. Możemy przyjąć, że jeżeli te trzy parametry będą poprawne, wzmacniacz ma zniekształcenia TIM na poziomie niewpływającym na jakość sygnału.

Na **rysunku 46** pokazano przykładowy pomiar zniekształceń TIM wzmacniacza ICEPower 250ASX250.

Tomasz Jabłoński, EP

REKLAMA



Kurs Nordic nRF z BT

Zanurzymy się w konfiguracji środowiska z nRF Connect SDK i przyjrzymy się, co sprawia, że płyta deweloperska nRF5340 DK jest tak wszechstronna. Przygotuj się na ekscytującą podróż przez konfigurację, programowanie oraz testowanie, które otworzą przed Tobą nowe możliwości w technologii Bluetooth Low Energy i systemie Zephyr.



ulubionykiosk.pl

Wejścia różnicowe w urządzeniach audio

W obecnych czasach prawie wszystkie akustyczne urządzenia elektroniczne budowane są przy użyciu cyfrowych układów przetwarzających sygnał dźwiękowy. Układy te łączą więc w sobie rozwiązania techniki analogowej i cyfrowej, m.in. przetworniki analogowo-cyfrowe (ADC) służące do przetwarzania sygnału akustycznego na postać binarną. Ponieważ do cyfrowej obróbki sygnału dźwiękowego potrzebna jest spora moc obliczeniowa, urządzenia te pobierają stosunkowo dużą moc z układu zasilania oraz pracują z wysoką częstotliwością, co z kolei sprawia, że emitują znaczące zaburzenia. W związku z tym analogowe układy wejściowe powinny być tak skonstruowane, aby dostarczyć sygnał do wejść przetwornika (ADC) z możliwie jak najmniejszymi zakłóceniami.

Korzyści ze stosowania wejścia różnicowego

W układach, w których skład wchodzi elementy cyfrowe i analogowe, do przesyłania sygnałów napięciowych najlepiej stosować linie różnicowe. Takie rozwiązanie ma szereg zalet: jest odporne na tętnienia napięcia zasilania, pasożytnicze spadki napięcia na połączeniach masy (**rysunek 1**) oraz – przy odpowiednim poprowadzeniu ścieżek – na zakłócenia elektromagnetyczne. Linie różnicowe zapewniają również dokładny pomiar składowej stałej napięcia wejściowego.

Ponieważ sygnały na obu liniach (1), (2) odejmują się, w efekcie otrzymujemy mierzone napięcie różnicowe praktycznie pozbawione zniekształceń (3):

$$U_{+C} = U_{+} + U_{noise} \quad (1)$$

$$U_{-C} = U_{-} + U_{noise} \quad (2)$$

$$U_r = U_{+C} - U_{-C} =$$

$$(U_{+} + U_{noise}) - (U_{-} + U_{noise}) = U_{+} - U_{-} \quad (3)$$

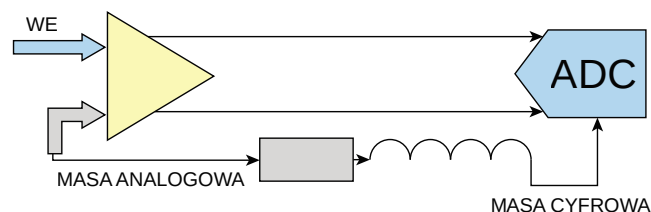
U_{+C} , U_{-C} wartości napięcia na liniach różnicowych razem z zakłóceniami,

U_r – napięcie różnicowe,

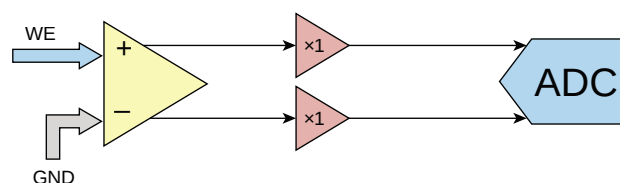
U_{noise} – składowa napięcia wejściowego reprezentująca zakłócenia.

Obwód wejściowy dla przetwornika ADC w mikrokontrolerze

Mikrokontroler przetwarzający sygnał akustyczny jest elementem, który wprowadza bardzo duże zakłócenia w układzie. Mają na to wpływ również obwody obecne w otoczeniu procesora. W związku z tym analogowy obwód wejściowy powinno się konstruować tak, aby był odporny na owe zaburzenia i dostarczał do przetwornika ADC procesora sygnał bez zakłóceń. Cel ten jest możliwy do osiągnięcia, jeśli wzmacniacz wejściowy będzie chroniony ekranem i wyposażony w oddzielną masę oraz wyjście różnicowe.



Rysunek 1. Wyjaśnienie drogi propagacji zakłóceń przez masę w układzie ze wzmacniaczem typu single-ended



Rysunek 2. Schemat blokowy stopnia wejściowego ze wzmacniaczem różnicowym

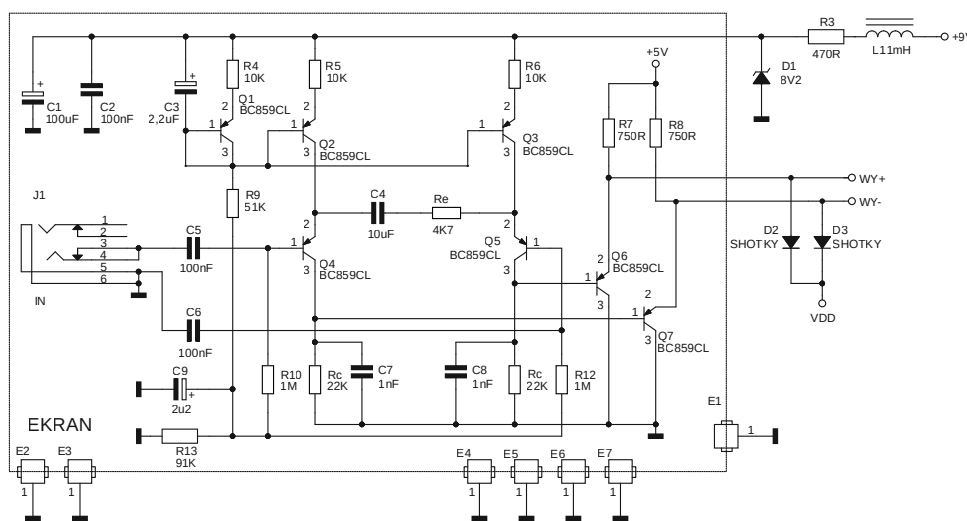
Ponadto – aby wykorzystać cały zakres napięć wejściowych przetwornika – powinno się zastosować taki układ, który dostarczy symetryczne wartości napięć o dokładnie ustalonym napięciu wspólnym, równym połowie zakresu napięć wejściowych.

Układ ze wzmacniaczem różnicowym

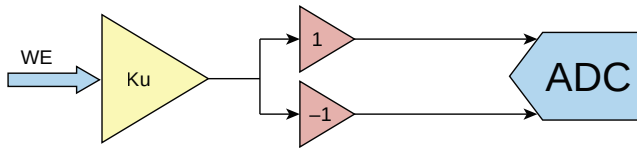
Wzmacniacz wejściowy można zaprojektować w postaci klasycznego wzmacniacza różnicowego. Układ ten charakteryzuje się prostotą konstrukcji, niską ceną i doskonałą odpornością na warunki pracy. Jest odporny na tętnienia napięcia zasilania, jak również na zakłócenia emitowane przez sam układ oraz środowisko zewnętrzne. Na wyjściu powinno się zastosować bufor – najlepiej w postaci wtórników emiterowych.

Schemat blokowy takiego rozwiązania pokazano na **rysunku 2**, natomiast na **rysunku 3** widać jego przykładową realizację praktyczną.

Układ ten można zmodyfikować, wyposażając go w przestrzajane wzmocnienie. Regulacja odbywa się poprzez zmianę wartości rezystora R_e np. przy zastosowaniu potencjometru cyfrowego. Wartość wzmocnienia opisuje wzór (4).



Rysunek 3. Schemat elektryczny stopnia wejściowego ze wzmacniaczem różnicowym



Rysunek 4. Schemat blokowy rozwiązania stopnia wejściowego z odwróceniem fazy

$$K_u = \frac{R_c}{\frac{2}{gm} + R_e} \quad (4)$$

Układ z obwodem odwracającym fazę

W przeciwieństwie do rozwiązania z poprzedniego punktu, w którym powstanie symetrycznego napięcia wynika z samej idei działania układu, w tym przypadku napięcie o przeciwnej fazie otrzymujemy poprzez zastosowanie wzmacniacza odwracającego. Ideę tego rozwiązania prezentuje **rysunek 4**.

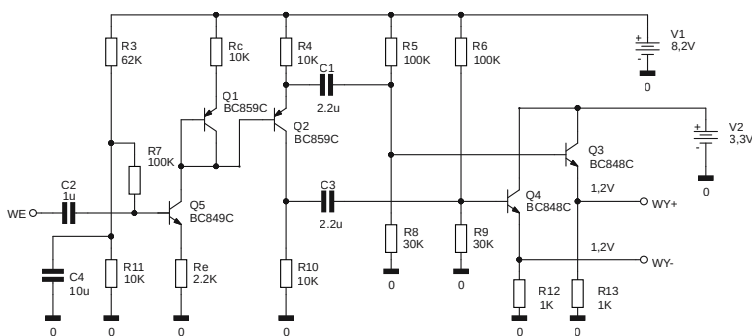
Wzmacniacz tranzystorowy

W układzie z **rysunku 5** zastosowano wzmacniacz napięciowy (zbudowany przy użyciu tranzystora pracującego w układzie wspólnego emitera) oraz układ odwracający fazę. Różnicowe sygnały wyjściowe wzmacniane są poprzez wtórnik emiterowy. W układzie tym również można regulować wzmocnienie, zmieniając wartość rezystora w układzie emitera tranzystora Q5 (5):

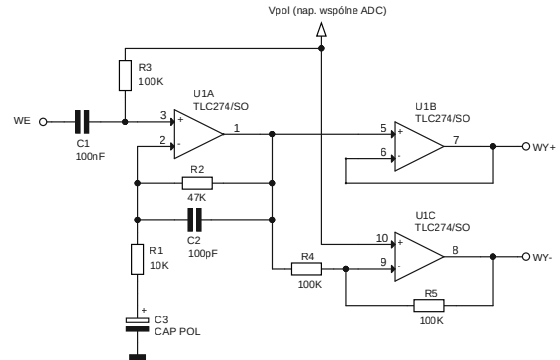
$$K_u = \frac{R_c}{\frac{1}{gm} + R_e} \quad (5)$$

Układ ze wzmacniaczami operacyjnymi

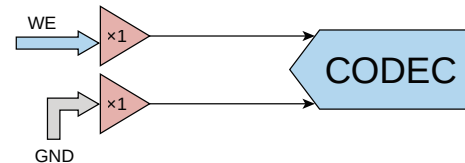
To rozwiązanie zawiera stopień wzmacniający w klasycznym układzie nieodwracającym wzmacniacza napięciowego ze sprzężeniem zwrotnym oraz wtórnik (odwracający i nieodwracający), pełniące również funkcję buforów wyjściowych. Zastosowanie wzmacniaczy operacyjnych sprawia, że układ okazuje się znacznie dokładniejszy, wprowadza mniejsze zniekształcenia nieliniowe, jednak jest droższy, wykazuje większe szумы oraz podatność na występowanie zniekształceń TIM. Przykładowy schemat takiego wzmacniacza pokazano



Rysunek 5. Schemat ideowy tranzystorowego wzmacniacza różnicowego



Rysunek 6. Schemat ideowy stopnia wejściowego ze wzmacniaczami operacyjnymi



Rysunek 7. Zastosowanie wtórników w celu zmniejszenia zakłóceń

na **rysunku 6**. O wartości wzmocnienia sygnałów zmiennych decyduje dzielnik złożony z rezystorów R1 i R2 (6):

$$K_u = 1 + \frac{R_2}{R_1} \quad (6)$$

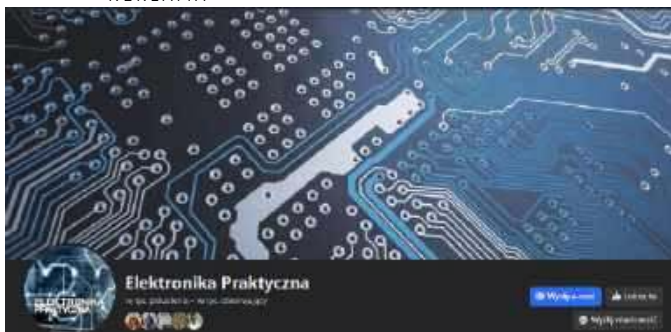
Obwód wejściowy dla przetwornika ADC w CODEC-u

Chociaż wysokiej klasy układy CODEC-ów mają dużą rezystancję wejściową, to wcale – jeśli chodzi o poziom zakłóceń – nie jest to cecha pożądana. Gdy konstruujemy własny, zaawansowany wzmacniacz wejściowy, CODEC również nie powinien mieć dużej czułości. Wniosek? Najlepiej stosować przetwornik o wejściu różnicowym – i nie jest tu wymagana duża rezystancja wejściowa czy czułość.

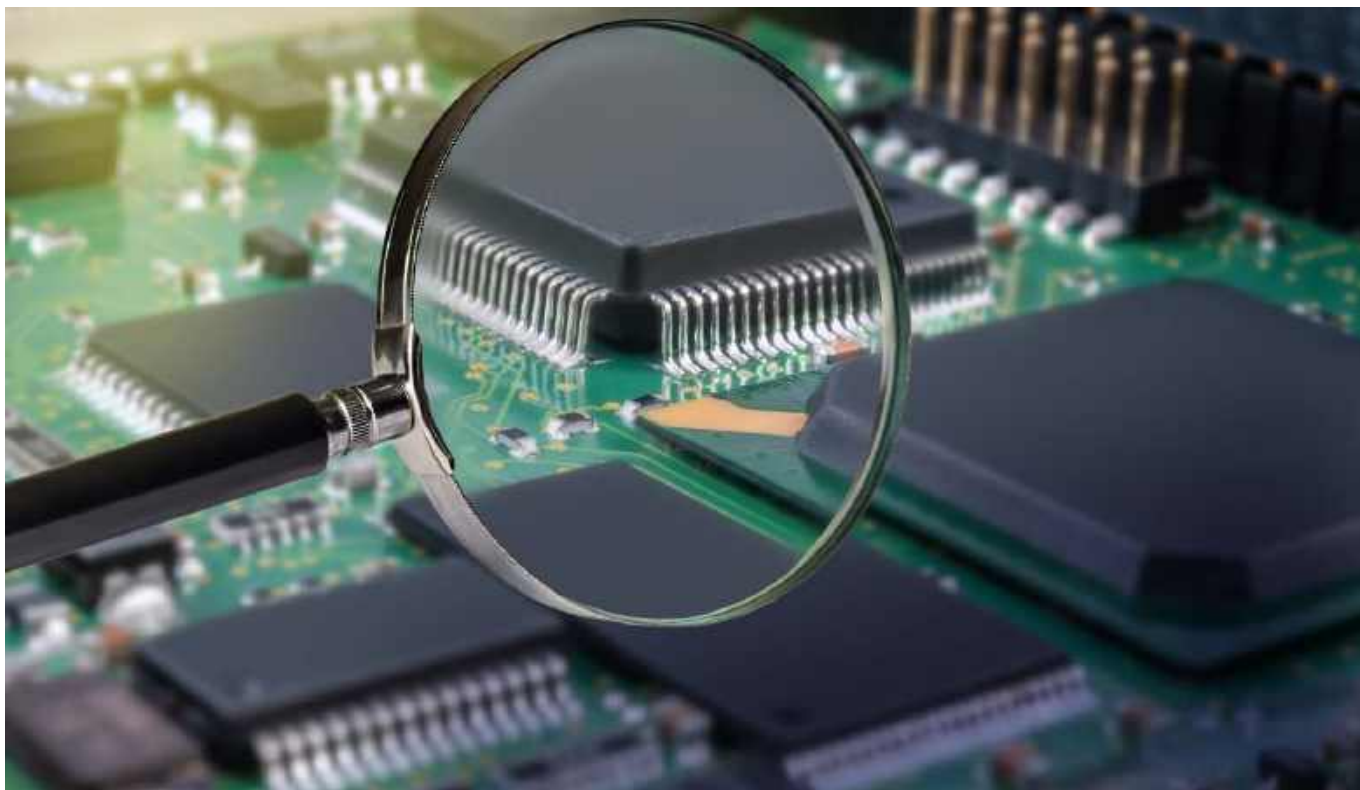
Najprostszym rozwiązaniem wejściowego obwodu CODEC-a okazuje się zastosowanie dwóch wtórników podłączonych do obydwu wejść (**rysunek 7**). Konfiguracja różnicowa zapewnia redukcję zakłóceń wynikających z różnic potencjałów pomiędzy masami obwodów cyfrowego i analogowego, natomiast niska rezystancja wyjściowa wtórników powoduje, że na ścieżkach sygnałowych nie indukują się zakłócenia elektromagnetyczne. Ponadto – jeśli poprowadzimy obie ścieżki sygnału analogowo blisko siebie – otrzymamy dodatkową redukcję zakłóceń. Natomiast jeżeli chcemy, aby sygnał był dodatkowo wzmocniony, możemy stosować układy omówione w poprzedniej części artykułu. Przy tych rozwiązaniach warto część analogową umieścić w ekranie, natomiast układ cyfrowy – poza nim, co dodatkowo sprawi, że wzmacniacz wejściowy będzie odporny na zakłócenia emitowane przez część cyfrową.

Tomasz Krogulski

REKLAMA



[www.facebook.com/
ElektronikaPraktyczna](http://www.facebook.com/ElektronikaPraktyczna)



Oszczędzanie energii w teorii i w praktyce (4)

W poprzedniej części niniejszego cyklu przyjrzelśmy się różnym czujnikom i układom z nimi współpracującym. Komponenty te często realizują specyficzne zadania, jak na przykład pomiar tętna i natlenienia krwi czy rozpoznawanie różnych ruchów i gestów. Jednak nie są to jedyne specjalizowane układy współpracujące z mikrokontrolerami.

W tym odcinku przyjrzymy się układom zastępującym i uzupełniającym wbudowane peryferia mikrokontrolera. Są to na przykład układy watchdog, generatory kwarcowe czy zegary czasu rzeczywistego. Przyjrzymy się też układom POR/BOR, które dbają o to, by procesor pracował tylko przy właściwym napięciu zasilania. Warto spojrzeć również na – często występujące w mikrokontrolerach – wzmacniacze operacyjne i sprawdzić, czy zewnętrzne układy scalone (lub wręcz rozwiązania tranzystorowe) nie okażą się lepsze. To ostatnie pytanie jest szczególnie ciekawe, zwłaszcza biorąc pod uwagę, jak wielkie mamy bogactwo układów scalonych i jak rzadko obecnie trafia się w komercyjnych urządzeniach na tranzystory realizujące inne funkcje niż przełączanie.

Układy nadzorcze: POR, BOR i watchdog

W praktyce każdy mikrokontroler ma wbudowane obwody POR, BOR i watchdog. Układ POR utrzymuje procesor w stanie resetu, póki napięcie zasilania nie osiągnie stabilnej wartości minimalnej. W przeciwnym razie program może być wykonywany niepoprawnie. Zakończenie stanu POR powoduje też ustawienie wartości domyślnych we wszystkich rejestrach (chyba że nota katalogowa mikrokontrolera podaje inaczej) oraz wyzerowanie pamięci RAM.



Poprzednie odcinki znajdują się pod adresem:
<https://ulubionykiosk.pl/media>

Z kolei układ BOR wstrzymuje pracę programu, gdy napięcie spadnie poniżej minimalnego progu. Reset spowodowany spadkiem napięcia zasilania jest zazwyczaj odnotowywany przez mikrokontroler za pomocą flagi w jednym z rejestrów, przy czym jej stan może zostać zachowany przez dłuższy czas po zaniku zasilania. Z praktycznego punktu widzenia funkcjonalność tę można zastosować do zmiany sposobu pracy układu przez szybkie jego wyłączenie i włączanie – tak działają latarki z wieloma trybami świecenia przełączanymi przyciskiem zasilania. Dlaczego zatem ktoś miałby decydować się na użycie zewnętrznego układu POR/BOR, skoro mikrokontroler już taki zawiera? To samo pytanie można zadać w odniesieniu do zewnętrznych układów watchdog, które – podobnie do wewnętrznego watchdoga procesora – mają za zadanie zresetować układ po upływie określonego czasu, jeśli program wcześniej nie da znaku życia watchdogowi. Moduł watchdog wewnątrz mikrokontrolera w wielu układach może wybudzić ten procesor ze stanu uśpienia, bez utraty kontekstu sprzed uśpienia – program kontynuuje pracę od miejsca, w którym uśpienie zostało zainicjowane. Zewnętrzny układ watchdog zresetuje mikrokontroler całkowicie, rozpoczynając egzekucję programu od początku. Czytelnik może zatem dojść do wniosku, że zewnętrzny watchdog jest mało użyteczny i niepotrzebnie podnosi koszty, podobnie zresztą jak układy POR/BOR.

Zewnętrzne układy POR/BOR i watchdog mają jednak kilka cech, które mogą być pożądane przez projektanta. Po pierwsze: zwykle są bardziej energooszczędne od modułów wbudowanych

w mikrokontroler, szczególnie gdy rzeczony procesor jest starszej generacji. Po drugie układy POR/BOR mają szeroki wybór napięć granicznych, dzięki czemu można precyzyjnie dobrać punkt wyłączenia w układzie zasilanym bateryjnie. Inną zaletą zewnętrznego układu POR/BOR może być precyzyjnie określone opóźnienie między ustabilizowaniem się napięcia zasilania a startem procesora. Większość mikrokontrolerów ma specjalny timer opóźniający start po osiągnięciu stosownego napięcia na pinie RESET – pin ten zazwyczaj dołączony jest do obwodu RC składającego się z rezystora 10 kΩ i kondensatora 100 nF. Jak łatwo się domyślić, czas aktywacji tak wyposażonej linii RESET zależy zatem nie tylko od stałej czasowej tego obwodu, ale także od napięcia zasilania. Zewnętrzny układ POR/BOR zazwyczaj ma stałe opóźnienie, które może zostać dobrane zależnie od potrzeb. Im mniej pewne źródło zasilania, tym dłuższe powinno być opóźnienie między ustabilizowaniem się napięcia a startem układu. Dla przykładu układ TPS3808 firmy Texas Instruments pozwala uzyskać opóźnienie wynoszące nawet dziesięć sekund. Pobór prądu tego układu – ze względu na jego wiek – jest znaczny i sięga typowo 2,4 μA. Z drugiej strony – to wciąż mniej niż wynosi pobór prądu modułu BOR w starszym mikrokontrolerze, na przykład PIC18LF45K50, tam bowiem najniższa typowa wartość wynosi 3,4 μA, a najwyższa typowa – aż 8 μA. W przypadku zaprezentowanego w części drugiej mikrokontrolera PIC18F27Q10 natężenie prądu zasilania BOR wynosi aż 30 μA, z kolei moduł LPBOR typowo pobiera już zaledwie 0,9 μA.

Spójrzmy zatem na kilka układów POR/BOR dostępnych na rynku. Na pierwszy ogień niech pójdzie seria TLV803E/809E/810E. Pobierają one typowo 250 nA, maksymalnie 1 μA. Różnica między nimi tkwi w sposobie działania wyjścia RESET – jest to odpowiednio: open-drain (TLV803E), aktywny stan niski push-pull (TLV809E), aktywny stan wysoki push-pull (TLV810E). Do wyboru mamy szereg napięć progowych: 1,7 V, 1,8 V, 1,9 V, 2,25 V, 2,4 V, 2,64 V, 2,93 V, 3,08 V, 3,3 V, 4,2 V, 4,38 V, 4,55 V, 4,63 V; i opóźnień: 40 μs, 10 ms, 20 ms, 50 ms, 100 ms, 200 ms i 400 ms. Histereza wynosi 1,2% napięcia przełączania, a dokładność – typowo 0,5%. Zależnie od obudowy dostępny może być również pin do dołączenia zewnętrznego przycisku resetu. Układy Maxim MAX6326/-7/-8 oraz MAX6346/-7/-8 oferują mniej opcji, jeśli idzie o napięcia przełączania, gdyż są przeznaczone do systemów o napięciu pracy 2,5 V, 3 V, 3,3 V i 5 V. Do wyboru jest szereg konfiguracji wyjścia, a pobór prądu w systemie 3 V wynosi tylko 500 nA typowo oraz 1 μA maksymalnie. Firma Torex z kolei oferuje kilka różnych serii układów nadzorujących zasilanie. Układy serii XC6126 oferują dość niski pobór prądu, 0,7 μA przy 3 V. Do wyboru są napięcia przełączania od 1,5 V do 5,5 V w krokach co 0,1 V. XC6129 oferuje podobne parametry, ale przy niższym poborze prądu, i dodaje możliwość wyboru opóźnienia resetu za pomocą zewnętrznego kondensatora (zwarcie tego kondensatora pozwala uzyskać ręczny reset).

Ciekawym wyborem jest też układ TLV840-Q1 firmy Texas Instruments, który oferuje wszystkie typowe możliwości układów POR/BOR, czyli wybór rodzaju wyjścia, wejście dla przycisku reset i szeroką gamę napięć progowych. Pobiera przy tym zaledwie 120 nA. Zewnętrzny kondensator pozwala wybrać opóźnienie, bez niego wynosi ono 40 μs, z kondensatorem 1 μF zaś 619 ms. Zakres napięć przełączania rozciąga się od 0,8 V do 5,4 V. Innymi układami o podobnym poborze prądu są te z rodziny MAX16056...MAX16059, które dodatkowo realizują funkcję watchdoga. W tym wypadku prąd zasilania wynosi 132 nA przy 3,3 V (typ.) – czyli o połowę mniej niż pobiera watchdog w układzie PIC18F27Q10. Wadą układu jest brak możliwości zmiany interwału watchdoga z poziomu oprogramowania, gdyż ta zależy od użytego kondensatora. Analog Devices w swojej ofercie ma dwa spokrewnione ze sobą układy: ADM696 i ADM697. Monitorują one napięcie zasilania, przełączają je z zewnętrznego źródła na zasilanie bateryjne i jednocześnie realizują funkcję watchdoga oraz POR/BOR. Przy zewnętrznym zasilaniu

pobierają do 1 mA, ale z baterii już maksymalnie 1 μA. Do ich możliwych aplikacji należy bateryjne podtrzymywanie pamięci RAM w systemie mikroprocesorowym, gdy zaniknie zewnętrzne zasilanie. Po jego powrocie mikroprocesor jest restartowany, a jednocześnie zachowane zostają dane w pamięci RAM. Inny układ łączący w sobie funkcję monitora napięcia i watchdoga to TPS3813-Q1 od Texas Instruments. W jego przypadku pobór prądu wynosi typowo 9 μA, czyli dość sporo. W zamian za to sam watchdog jest typu okienkowego, tj. sygnał resetu musi pojawić się przed upływem maksymalnego czasu, ale po upływie czasu minimalnego – inaczej nastąpi reset. Rozwiązanie to przydaje się, gdy mikrokontroler znajduje się przypadkiem w pętli, która nie robi nic poza resetowaniem watchdoga. TPS3430 tej samej firmy to programowalny watchdog okienkowy kontrolowany za pomocą trzech pinów ustalających interwał i szerokość okienka. Układ typowo pobiera 10 μA. Układy Maxim MAX6821...MAX6825 również łączą w sobie funkcję POR/BOR, Watchdog i ręcznego resetu. Interwał watchdoga wynosi 1,6 s, pobór prądu zaś typowo 7 μA przy 3,6 V. Różne modele łączą różne kombinacje wyjść i wejść, każdy zaś oferuje dziewięć różnych napięć przełączania do wyboru, od 1,58 V do 4,63 V.

Jak widać, wybór układów jest dość duży, ale nie każda z dostępnych opcji jest ekstremalnie energooszczędna. Timery watchdog wymagają bowiem oscylatorów, które znacząco podnoszą pobór prądu. Natomiast układy POR/BOR realizują opóźnienia za pomocą kondensatorów ładowanych źródłem napięcia odniesienia przez wewnętrzny rezystor, dzięki czemu stała czasowa obwodu jest stabilna i niezależna od napięcia zasilania. Czytelnik może też rozważyć budowę własnego układu POR/BOR opartego na tranzystorach. Jest to jak najbardziej akceptowalne rozwiązanie, pod warunkiem, że cały obwód zmieści się na płytce drukowanej, będzie odpowiednio energooszczędny i bezawaryjny, jednocześnie oferując niższy koszt niż użycie specjalnego układu scalonego. Zaprojektowanie takiego rozwiązania zostawiam jako wyzwanie dla zainteresowanych Czytelników.

Generatory kwarcowe i układy RTCC

Niemal każdy mikrokontroler może pracować z zewnętrznym zegarem, zazwyczaj pobierając przy tym mniej prądu niż przy użyciu nawet najbardziej oszczędnej opcji wbudowanej. Użycie zewnętrznego generatora może mieć zatem całkiem sporo sensu, zwłaszcza jeśli jest to generator programowalny. Z reguły wewnętrzny oscylator mikrokontrolera ma tylko kilka do kilkunastu stałych częstotliwości, najczęściej każda jest 2 razy większa od poprzedniej. Jeśli zatem potrzebujemy częstotliwości niestandardowej, to jedyną opcją staje się zewnętrzny rezonator lub generator kwarcowy. Noty katalogowe nie podają zwykle dolnego limitu częstotliwości taktowania, więc w teorii układ może pracować z dowolnie niską częstotliwością, nawet rzędu pojedynczych herców, co ma też wpływ na pobór prądu. W celu miarodajnego porównania, w tabelach 1 i 2 zestawiono

Tabela 1. Typowe wartości poboru prądu mikrokontrolera PIC12LF1501

Model mikrokontrolera	PIC12LF1501	
Napięcie zasilania	1,8 V	3 V
ECM 32 kHz, Low Power	6 μA	8 μA
ECM 500 kHz, Low Power	19 μA	32 μA
ECM 1 MHz	30 μA	55 μA
ECM 4 MHz	115 μA	210 μA
ECH 20 MHz	–	1050 μA
LFINTOSC 31 kHz	3,2 μA	5,4 μA
HFINTOSC 500 kHz	215 μA	275 μA
HFINTOSC 8 MHz	410 μA	630 μA
HFINTOSC 16 MHz	600 μA	970 μA

Tabela 2. Typowe wartości poboru prądu mikrokontrolera PIC18LF45K50

Model mikrokontrolera	PIC18LF45K50	
Napięcie zasilania	1,8 V	3 V
ECM 1 MHz (PRI IDLE)	30 μ A	45 μ A
ECM 20 MHz (PRI IDLE)	450 μ A	700 μ A
ECM 1 MHz (PRI RUN)	110 μ A	170 μ A
ECH 20 MHz (PRI RUN)	1450 μ A	2650 μ A
ECM 4 MHz (4 $\times$ PLL – 16 MHz Fosc)	350 μ A	550 μ A
ECH 16 MHz (3 $\times$ PLL – 48 MHz Fosc)	–	6350 μ A
SOSC 32 kHz	1 μ A	1,4 μ A
HFINTOSC 1 MHz (RC IDLE)	250 μ A	350 μ A
HFINTOSC 16 MHz (RC IDLE)	500 μ A	800 μ A
HFINTOSC + 3 $\times$ PLL 48 MHz (RC IDLE)	–	2500 μ A

typowe pobory prądu dwóch mikrokontrolerów: PIC12LF1501 oraz PIC18LF45K50 przy różnych częstotliwościach i typach oscylatorów.

Jak łatwo zauważyć, pobór prądu przy zewnętrznym zegarze bywa znacząco niższy od prądu zasilania podczas pracy z użyciem wewnętrznego oscylatora. Zwraca też uwagę bogactwo opcji układu PIC18LF45K50 – ma on bowiem wbudowany moduł USB 2.0 wymagający zegara 48 MHz. Co ciekawe, pobór prądu wewnętrznych oscylatorów jest wciąż dość niski – znalezienie bardziej energooszczędnego generatora to pewne wyzwanie, zwłaszcza że do jego poboru prądu trzeba doliczyć też pobór prądu samego mikrokontrolera z zewnętrznym zegarem. W zestawieniu uwzględniono tylko kilka wybranych układów, jako że zdecydowana większość scalonych generatorów kwarcowych pobiera od kilku do kilkunastu miliamperów, daleko wykraczając poza możliwości samych mikrokontrolerów.

Najbardziej prądożerny układ to 5X1503 firmy Renesas. Oferuje on trzy niezależne wyjścia zegarowe, każde o częstotliwości od 1 MHz do aż 100 MHz – które mogą też dostarczać częstotliwość 32,768 kHz, co pozwala na współpracę z układami i modułami RTCC. Każde z wyjść ma wejście kontrolne, umożliwiające jego włączenie lub wyłączenie – da się zatem zaprogramować trzy wyjścia na różne częstotliwości, podłączyć do jednego wejścia zegarowego mikrokontrolera i używać wejść kontrolnych, by wybierać, która częstotliwość będzie używana. Układ 5X1503L pobiera przy tym 2 mA w czasie pracy i około 1 μ A w trybie oszczędnym z wyłączonym interfejsem I²C. Jedną z ciekawych funkcji układu stanowi możliwość zapisania czterech konfiguracji dzielnika częstotliwości dla PLL w pamięci nieulotnej OTP – pozwala to na natychmiastowy start generatora, bez konieczności wstępnej konfiguracji.

Firma Maxim Integrated ma w ofercie szereg układów zegarowych. Jednym z nich jest DS1090, generator kwarcowy o częstotliwości wyjściowej programowanej pojedynczym rezystorem. Rezystor ten ustala częstotliwość oscylacji między 4 MHz a 8 MHz. Następnie ta częstotliwość jest dzielona przez 1, 2, 4, 8, 16 lub 32 – zależnie od wybranego układu – po czym do sygnału zegarowego dodawany jest dither o wybieranych przez użytkownika parametrach, dzięki czemu generowany zegar staje się celowo nieznacznie rozstrojony, co redukuje zakłócenia elektromagnetyczne emitowane przez urządzenie. Wadę układu stanowi wąski zakres napięć zasilania: 3...5,5 V. Przy napięciu 3,3 V układ pobiera typowo 1,4 mA. Ponieważ częstotliwość pracy ustalana jest rezystorem włączonym między wejście układu a masę, można rozważyć użycie większej liczby rezystorów. Dla przykładu: łącząc szeregowo rezystory 15 k Ω i 75 k Ω , a następnie bocznikując ten drugi rezystor tranzystorem, można przełączyć układ między skrajnymi częstotliwościami pracy. Ciekawsze rozwiązanie stanowi układ DS1099, który ma wewnętrzny oscylator 1,048 MHz i dwa niezależne dzielniki częstotliwości z wejściami sterującymi. Układ zasilany jest napięciem 2,7...5,5 V, pobierając w stanie spoczynku tylko

140 μ A. Częstotliwość wyjściowa zależy od wartości dzielnika 2n, gdzie n ma wartość od 0 do 22. Wartości podziałów są programowane w fabryce i trzeba je określić przy składaniu zamówienia. Przy maksymalnej wybranej częstotliwości w stanie aktywnym układ pobiera 323 μ A. Dodatkowo wydajność prądowa wyjść jest wystarczająca do bezpośredniego podłączenia diody LED – dlatego też minimalna dostępna częstotliwość wynosi $\sim$ 0,25 Hz.

Czytelnik zapewne zorientował się już, że nawet najbardziej oszczędny układ nie jest aż tak oszczędny, jak wewnętrzny oscylator mikrokontrolera. Tylko układ firmy Renesas – w połączeniu z mikrokontrolerem PIC18LF45K50 – może zyskać kilkaset mikroamperów oszczędności dla częstotliwości od 20 MHz. Z drugiej strony, z układem DS1099 można uzyskać częstotliwości zegarowe dużo niższe, niż pozwala na to sam mikrokontroler, co przy braku dolnego limitu zegara i zależności poboru prądu od częstotliwości taktowania może uczynić takie rozwiązanie bardziej energooszczędnym. W takich sytuacjach zarówno DS1099 firmy Maxim, jak i LTC6991 od Analog Devices mają więcej sensu. Ten drugi układ pozwala na generowanie sygnału prostokątnego o okresie od 1 ms do 9,5 godziny, więc może też pełnić funkcję watchdog. Pobór prądu sięga od 55 do 80 μ A – zależnie od wybranego okresu – ale napięcie zasilania to minimum 2,25 V. Czytelnikom, którzy mogą niechętnie patrzeć na częstotliwość taktowania do 1 kHz, pragnę przypomnieć, że wczesne komputery oraz wiele elektronicznych kalkulatorów miało podobne lub niższe częstotliwości pracy.

Jeśli tworzoną projektem jest układ realizujący funkcję zegarka (lub na potrzeby gromadzenia danych niezbędne są data i godzina), to potrzebny może okazać się układ RTCC, czyli zegar czasu rzeczywistego. Część mikrokontrolerów ma taki moduł już wbudowany, reszta wymaga zewnętrznego układu. Rozwiązania te można podzielić na dwie kategorie: wyposażone we wbudowany oscylator – lub wymagające zewnętrznego rezonatora bądź generatora kwarcowego, zazwyczaj 32,768 kHz. Warto tu zaznaczyć, że większość rezonatorów zegarkowych jest zoptymalizowana do pracy w temperaturze zbliżonej do temperatury ludzkiego ciała, dlatego poniżej lub powyżej tej temperatury częstotliwość rezonansowa zaczyna spadać. Drugą powszechnie występującą funkcjonalnością układów zegara czasu rzeczywistego są przerwania/alarmy. Wzorcowym przedstawicielem układów RTCC jest MAX3131 firmy Analog Devices. Układ ten wymaga do pracy jedynie rezonatora kwarcowego oraz rezystorów podciągających, jako że wszystkie jego wyjścia są typu otwarty dren. Kontrolowany jest za pomocą interfejsu I²C oraz jednego dodatkowego pinu sterującego; ma też dwa wyjścia przerwań, przy czym jedno z nich może być skonfigurowane jako wyjście zegarowe. Układ pracuje z napięciem od 1,1 V do 5,5 V, ale po spadku napięcia zasilania poniżej ustawionego progu (1,55 V albo 2 V) przełącza się w tryb zasilania z baterii lub superkondensatora – w tym trybie interfejs komunikacyjny jest wyłączony, a układ pobiera typowo 65 nA przy 3 V na potrzeby podtrzymania funkcji zegara. Gdy zasilanie główne jest włączone, pobór rośnie ze względu na prąd upływu wejść i wyjść. Dodatkowo może też działać funkcja doładowania baterii niewielkim prądem przez jeden z trzech rezystorów wewnątrz układu: 3,3 k Ω , 6,4 k Ω i 11,3 k Ω . Poza standardowymi funkcjami zegara z alarmem układ ma też tryb timera z pauzą oraz możliwość zachowania dokładnej daty i godziny zdarzenia (wygenerowanego przez pin wejściowy – lub zaniku i powrotu głównego zasilania).

Jeszcze ciekawiej wygląda rodzina układów AM16x5 Artasie firmy Ambiq. Zliczają one czas od ułamków sekundy, mają też wewnętrzny oscylator RC oraz współpracują z zewnętrznym rezonatorem zegarkowym. Rozwiązanie to pozwala na automatyczną kalibrację oscylatora RC i użycie go celem zaoszczędzenia energii, podczas gdy obwód z rezonatorem kwarcowym jest używany sporadycznie (w momencie podjęcia takiej decyzji przez mikrokontroler). Napięcie zasilania mieści się w zakresie od 1,5 V do 3,6 V, pobór

prądu zaś zależy od trybu pracy: z rezonatorem kwarcowym przy 3 V układ pobiera 55 nA, z oscylatorem RC – już tylko 14 nA. W trybie autokalibracji, co 512 sekund (oscylator kwarcowy jest włączany tylko na czas kalibracji) średni pobór prądu wynosi 22 nA. Przy zasilaniu z wejścia baterijnego prądy te są o 1...2 nA wyższe. Układ wspiera też ładowanie baterii, używając wewnętrznych rezystorów o wartościach 3 kΩ, 6 kΩ i 11 kΩ. Seria AM16x5 oferuje dwa interfejsy do wyboru: I²C i SPI. W czasie przesyłania danych przez interfejs I²C układ pobiera typowo 10 μA, natomiast w przypadku interfejsu SPI jest to 8 μA przy zapisie do układu i 23 μA przy odczycie z układu. Jedną z potencjalnie przydatnych funkcji stanowi generator przebiegów prostokątnych, którego częstotliwość jest wybierana w zakresie od 32,768 kHz po nawet... jeden wiek. Oczywiście niektóre częstotliwości dostępne są tylko z aktywnym zewnętrznym rezonatorem, zaś te powiązane z rejestrami daty i godziny służą głównie do celów testowych. Niemniej jednak to wyjście może być źródłem zegara mikrokontrolera albo przerwania. Wbudowany watchdog ma 5-bitowy licznik, dekrementowany z częstotliwością ¼ Hz, 1 Hz, 4 Hz lub 16 Hz, dzięki czemu uzyskać można okres watchdoga od 62,5 ms do 128 s. Układ ma specjalne wejście resetu watchdoga, a także rozbudowany system alarmów (wliczając w to alarm dla wybranych dni tygodnia) oraz wbudowany timer. W przeciwieństwie do MAX31331 nie oferuje funkcji zachowywania daty zdarzenia – mikrokontroler musi odczytać bieżącą datę i godzinę samodzielnie oraz zapisać je do pamięci.

Trzecim, przykładowym zegarem czasu rzeczywistego jest układ PCF85263A firmy NXP Semiconductors. Pod wieloma względami podobny do AM18x5, oferuje funkcję zapamiętywania daty i godziny zdarzenia. W trakcie pracy pobiera nieco więcej prądu, bo aż 280 nA przy 3 V, ale w zamian zakres napięć zasilania sięga w jego przypadku od 0,9 V do 5,5 V. Układ może zainteresować hobbystów, gdyż oferuje łatwiejsze w użyciu obudowy SOIC i TSSOP – wcześniej wymienione układy występują w trudniejszych do lutowania obudowach QFN i WLP. W przeciwieństwie do powyżej wspomnianych układów, w tym zegarze nie ma funkcji timera – zastąpił ją stoper. Nie ma też wewnętrznego oscylatora RC ani układu doładowywania baterii. Funkcja watchdoga jest identyczna w działaniu z AM18x5. Wyjście zegarowe też występuje, ale ograniczone jest do 32,768 kHz, 16,384 kHz, 8,192 kHz, 4,096 kHz, 2,048 kHz, 1,024 kHz i 1 Hz.

Wymienione układy RTCC reprezentują jedynie niewielki wycinek bogactwa ofert różnych producentów. Nie tylko są bardzo energooszczędne, ale też oferują wiele różnorodnych, użytecznych funkcji. Osiągnięcie tak niskich poborów prądu, często sporo poniżej mikroampera, było możliwe dzięki zamierzającemu już rynkowi zwykłych, elektronicznych zegarków naręcznych. Dążenie do osiągnięcia maksymalnej dokładności i minimalnego poboru prądu doprowadziło nie tylko do optymalizacji samych rezonatorów kwarcowych, ale też współpracujących z nimi dzielników i liczników. Pobór prądu przez te układy jest tak niski, że czas pracy układu RTCC ogranicza praktycznie tylko trwałość baterii i jej elektrolitu.

Wzmacniacze operacyjne kontra wzmacniacze tranzystorowe

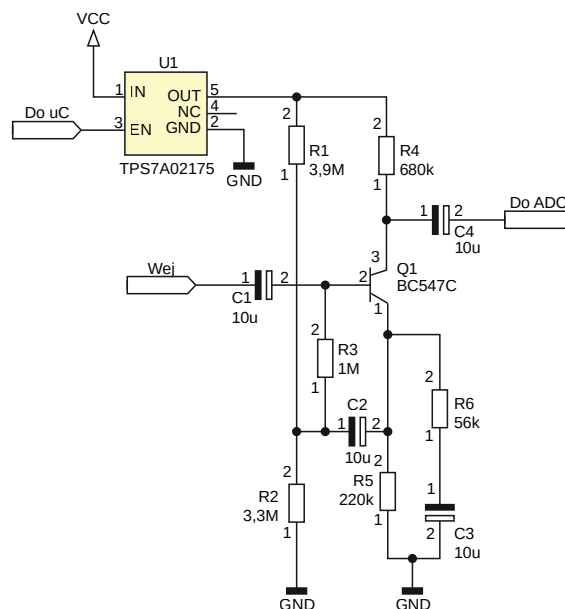
Jeśli projektowany układ dokonuje pomiaru napięć za pomocą ADC, konieczne może okazać się wzmocnienie mierzonego sygnału, by korzystać w pełni z rozdzielczości przetwornika. Niektóre mikrokontrolery mają wbudowane wzmacniacze operacyjne, w innych przypadkach może istnieć zapotrzebowanie na zewnętrzny układ. Oba te rozwiązania generują jeden problem: zwiększony pobór prądu. Tanie, popularne wzmacniacze operacyjne pobierają wręcz miliampery, nawet w stanie spoczynku. Układy określane jako „Ultra-low power”, czyli bardzo małej mocy, rzadko potrzebują mniej niż kilkudziesięciu mikroamperów na wzmacniacz. Dodatkowo sporo układów wymaga napięcia ponad 2,5 V,

podczas gdy sam mikrokontroler będzie pracować przy napięciu 1,8 V. Jeśli mamy miejsce, warto rozważyć zaprojektowanie własnego stopnia wzmacniającego, gdyż to rozwiązanie ma szansę okazać się zwyczajnie tańsze i bardziej energooszczędne.

Rozważmy więc hipotetyczną sytuację, gdy potrzebne jest wzmocnienie 10×, czyli 20 dB, pasmo przenoszenia do 20 kHz, jak najmniejszy pobór prądu oraz możliwość wyłączenia układu wzmacniacza, gdy mikrokontroler przechodzi w stan uśpienia. Wzmacniacz operacyjny powinien zatem mieć pasmo przenoszenia (GBW) powyżej 200 kHz [1] i być zoptymalizowany do zasilania pojedynczym napięciem od 1,8 V – takie cechy mają wzmacniacze rail-to-rail. Poza porównaniem kilku ofert gotowych układów sprawdzimy też rozwiązanie dyskretnie w symulacji.

Na pierwszy ogień idzie oferta Microchipa i dwie rodziny układów: MCP6001/1R/1U/2/4 oraz MCP6231/1R/1U/2/4. Pierwsza rodzina oferuje pasmo przenoszenia 1 MHz, ale energooszczędność wypada dużo słabiej – układy pobierają 100 μA na każdy kanał. Rodzina MCP623x prezentuje się lepiej pod tym względem: tylko 20 μA na kanał z pasmem 300 kHz. Kolejna rodzina tego producenta, MCP6006/6R/6U/7/9, ma pasmo do 1 MHz, pobór prądu 70 μA typowo, ale szczególną cechą tych układów jest podwyższona odporność na zakłócenia elektromagnetyczne wysokiej częstotliwości. Firma STMicroelectronics oferuje szereg różnych wzmacniaczy operacyjnych, na przykład serię TS185x. Układy te oferują pasmo do 530 kHz, wysoką stabilność nawet przy dużych obciążeniach pojemnościowych, ale i spory pobór prądu – aż 120 μA. Za to układy z rodziny TSV619x mogą się poszczycić poborem bardzo niskim, ledwo 10 μA na kanał, mogą być zasilane napięciem 1,5 V, pasmo sięga 450 kHz, jednak układy nie są stabilne jako bufory. Szersze pasmo, bo do 2,4 MHz, oferuje model TSV639x – niestety pobór prądu rośnie do 60 μA, a ponadto wzmacniacz nie jest stabilny jako bufor lub układ o małym wzmocnieniu. Oferuje za to wejście wyłączające, które ogranicza pobór prądu do znakomitej wartości 5 nA. Z oferty Texas Instruments najbardziej warta uwagi jest rodzina OPAx333. Układy te to precyzyjne wzmacniacze operacyjne o bardzo niskim dryfcie, niskich wartościach napięcia niezrównoważenia, szumów własnych i poboru prądu (tylko 17 μA), ich pasmo zaś wynosi 350 kHz.

Rozwiązanie dyskretnie konstrukcyjnie jest dość proste – składa się z pojedynczego tranzystora i garści innych komponentów. Proste układy tranzystorowe mają jednak dość poważną wadę: dużą zależność wzmocnienia od napięcia zasilania. Korekta tego problemu



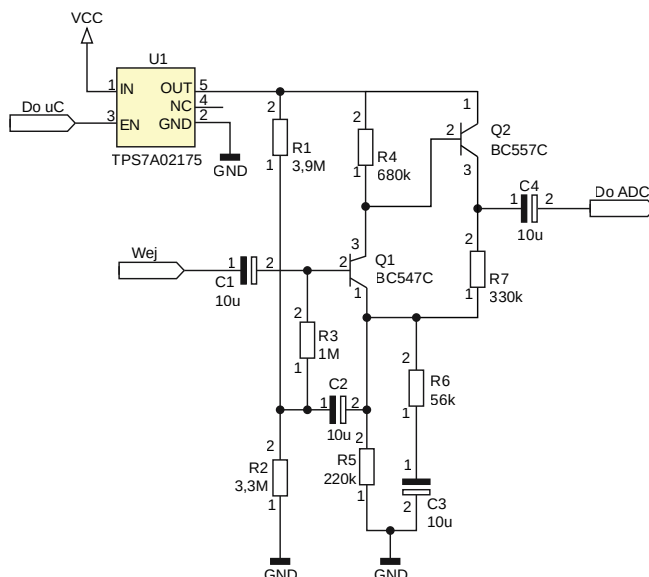
Rysunek 1. Prosty, energooszczędny wzmacniacz na pojedynczym tranzystorze z funkcją wyłączania przez mikrokontroler

wymaga rozbudowania układu do poziomu przypominającego schemat wewnętrzny wzmacniacza operacyjnego, co byłoby nielogiczne. Zatem – by rozwiązanie tranzystorowe miało sens – potrzebne jest stabilne źródło napięcia o możliwie niskim własnym poborze prądu. Takie warunki spełnia układ TPS7A02 firmy Texas Instruments. Pobiera on 25 nA w trakcie pracy, 3 nA, gdy układ jest wyłączony.

Rysunek 1 pokazuje praktyczny schemat z układem TPS7A02175: stopień wzmacniający zasilany jest napięciem 1,75 V, wzmocnienie wynosi niemal 20 dB, a pasmo sięga do ok. 40 kHz, pobór prądu zaś wynosi 1,66 μ A. Czytelnik zauważy, że wartości rezystorów są znacznie wyższe niż zazwyczaj spotykane w takich układach. Właśnie dzięki temu ograniczono pobór prądu kosztem wyższej impedancji wyjściowej. By przy okazji poprawić impedancję wejściową, zastosowano bootstrap składający się z R3 i C2. Amplituda sygnału wejściowego jest ograniczona przez napięcie zasilania do około 80 mV p-p, na wyjściu wyniesie zatem 800 mV p-p. Przetworniki ADC mikrokontrolerów z reguły potrzebują źródła sygnału o impedancji do 10 k Ω , ale to prawda przy założeniu maksymalnej częstotliwości próbkowania – źródło musi mieć na tyle wysoką wydajność prądową, by zdążyć przeładować kondensator próbkujący o wartości kilku do kilkunastu pikofaradów. **Rysunek 2** pokazuje zmodyfikowany wariant układu z dodatkowym tranzystorem. Pobór prądu rośnie do około 3,5 μ A, ale amplituda sygnału wyjściowego może już osiągnąć wartość 1 V p-p przy jednoczesnym obniżeniu impedancji wyjściowej. Pasmo przenoszenia sięga 30 kHz.

Zakończenie

Artykuł ten nie wyczerpuje listy alternatywnych układów peryferyjnych dla mikrokontrolerów, a w wymienionych kategoriach ogranicza się tylko do kilku przedstawicieli wybranych na podstawie ich parametrów. Najciekawiej wypadła kwestia wzmacniaczy operacyjnych, które pod względem poboru prądu przegrywają z prostymi układami tranzystorowymi, ale zyskują przewagę, gdy porównujemy inne parametry, takie jak pasmo przenoszenia czy niezależność wzmocnienia od napięcia pracy. W następnej części cyklu przyjrzymy się układom i modułom do komunikacji radiowej, a także zmierzmy się



Rysunek 2. Rozbudowany wariant wzmacniacza, poprawiający jego parametry kosztem nieco zwiększonego poboru prądu

z problemem dużego poboru prądu w czasie nadawania/odbioru oraz poszukamy sposobu na oszczędzanie energii nawet w czasie nasłuchu.

Paweł Kowalczyk, EP

Przypisy:

[1] Popularna reguła inżynierska sugeruje obliczenie minimalnego GBW wzmacniacza operacyjnego jako iloczynu częstotliwości maksymalnej, wzmocnienia w zamkniętej pętli sprzężenia zwrotnego oraz arbitralnie przyjętego współczynnika bezpieczeństwa równego przynajmniej 50...100, w celu optymalizacji parametrów układu. W omawianym przypadku byłoby to 10...20 MHz – zastosowanie mniejszego zapasu wzmocnienia pozwala obniżyć koszty (a zwykle także pobór mocy), ale pogarsza inne parametry sygnałowe układu – przyp. red.

REKLAMA

UWAGA! Tylko prenumeratorzy czasopism „Elektronika dla Wszystkich”, „Elektronika Praktyczna”, „Świat Radio” oraz „Elektronik” mogą korzystać z atrakcyjnych rabatów w Sklepie AVT:

- ✓ do 50% na wydania specjalne czasopism Wydawnictwa AVT
- ✓ 20% na kity w wersji A (płytki drukowane do projektów AVT)
- ✓ 10% na pozostałe wersje kitów: (A+, B, C, D)
- ✓ 10% na książki
- ✓ 5% na pozostałe produkty z oferty sklepu

Ponadto każdy prenumerator ww. czasopism korzysta z rabatów od 30% do 50% na zakup czasopism z oferty www.UlubionyKiosk.pl

K L U B
AVT
ELEKTRONIKA

Jak uzyskać rabat? Podczas zamówienia powołaj się na swój numer prenumeraty – otrzymasz go mailowo po zakupie prenumeraty wraz z kartą członkowską Klubu AVT-Elektronika.

Regulamin Klubu AVT-Elektronika znajdziesz na stronie <https://sklep.avt.pl/klub-avt-elektronika>

Kurs FPGA Lattice (24)

Miernik częstotliwości

W tym odcinku kursu nie poznamy niczego nowego. Będziemy natomiast ćwiczyć w praktyce stosowanie modułów opracowanych we wcześniejszych odcinkach. Zbudujemy miernik częstotliwości, w którym wynik pomiaru prezentowany będzie na wyświetlaczu 7-segmentowym.

Częstotliwość, według definicji, jest to wielkość fizyczna określająca liczbę wystąpień jakiegoś zjawiska w jednostce czasu. Jednostką czasu w układzie SI jest sekunda, a zjawiskiem, jakie chcemy badać, jest zbrocze rosnące mierzonego sygnału prostokątnego. Zatem: aby zmierzyć częstotliwość sygnału, musimy po prostu policzyć, ile występuje zbroczy tego sygnału w ciągu jednej sekundy.

Na płytce User Interface Board nie mamy żadnych układów umożliwiających kondycjonowanie różnych sygnałów – badany przebieg doprowadzony będzie prosto do pinu układu FPGA. Wynika z tego bardzo ważne ograniczenie: sygnał pochodzący z zewnętrznego generatora musi mieścić się w przedziale od 0 do 3,3 V. Badany przebieg powinien mieć kształt prostokątny – pomiar fali sinusoidalnej, trójkątnej lub innej może dać niepoprawny wynik.

Przeanalizujemy schemat z **rysunku 1**, wygenerowany automatycznie przez narzędzie Netlist Analyzer na podstawie syntezy kodu źródłowego w Verilogu, aby lepiej zrozumieć ideę działania miernika częstotliwości.

Układ będzie miał trzy wejścia. Do wejścia zegarowego **Clock** podłączony zostanie generator kwarcowy o częstotliwości 25 MHz, umieszczony na płycie MachXO2 Mega, prezentowanej już w EP 09/2023. Wejście **Reset** działa dokładnie tak samo, jak we wszystkich dotychczasowych projektach. Do wejścia **SignalAsync_i** doprowadzić należy sygnał, którego częstotliwość ma zostać zmierzona.

Badany przebieg oczywiście nie jest zsynchronizowany z domeną zegarową. Aby uniknąć problemu metastabilności, należy go najpierw zsynchronizować. W tym celu trzeba ów sygnał „przepuścić” przez moduł **Synchronizer**, widoczny po lewej stronie schematu. Temat ten został bardzo obszernie omówiony w 11 odcinku kursu poświęconym tematowi statycznej analizy czasowej (opublikowanym w EP 09/2023).

Synchronizator daje nam sygnał o takiej samej częstotliwości i wypełnieniu, jak sygnał wejściowy, ale my potrzebujemy wykrywać zbocza rosnące sygnału. W tym celu zastosujemy moduł **EdgeDetector**. Obserwuje on wejście **Signal_i**, a kiedy zmienia się jego stan, wówczas na wyjściach **RisingEdge_o** lub **FallingEdge_o** ustawiany jest – na jeden takt zegarowy – stan wysoki (w zależności od tego, czy moduł wykrył, odpowiednio, zobacze rosnące czy opadające). Nas interesuje tylko wykrywanie zbocza rosnącego i każde takie zdarzenie będziemy zliczać za pomocą licznika **Counter**.

Popatrzmy teraz na licznik **Counter** oraz jego wyjście **q**. Jest ono doprowadzone do wejścia **a** sumatora **add 5** (nazwa



Poprzednie odcinki znajdują się pod adresem:
<https://ulubionykiosk.pl/media>

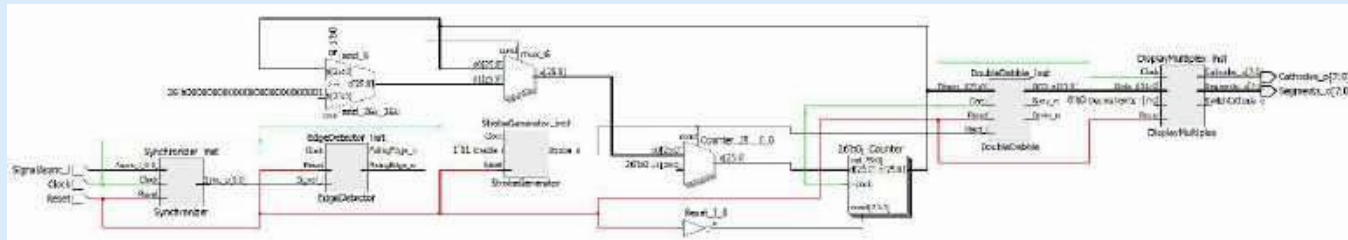
Zobacz więcej:

- Repozytorium modułów użytych w kursie:
<https://github.com/leonow32/verilog-fpga>
- Projekt w programie Diamond: <https://tiny.pl/9pkw6h-v>
- Film pokazujący działanie układu:
<https://youtu.be/FtbOKaUscUE>

wygenerowana została automatycznie), który do obecnego stanu licznika dodaje jedynkę, doprowadzoną „na sztywno” do wejścia **b**. Wyjście wspomnianego sumatora prowadzi do wejścia **d1** multiplexera **mux_6** (nazwa również wygenerowana automatycznie), a do wejścia **d0** doprowadzany jest aktualny stan licznika, niepowiększony o 1. Zatem multiplexer przepuści dalej aktualny stan lub stan powiększony o 1 – decyduje o tym wejście **cond** multiplexera, sterowane przez wyjście **RisingEdge_o** detektora zbocza. Kiedy zostanie wykryte zbocze rosnące, wówczas stan tego wejścia jest wysoki przez czas jednego taktu zegarowego, a multiplexer przekazuje do swojego wyjścia stan wejścia **d1**, czyli stan licznika zwiększony o 1. W stanie spoczynku, kiedy nic nie jest wykrywane, multiplexer łączy wyjście z wejściem **d0**, czyli przekazuje dalej obecny stan licznika.

Miernik częstotliwości potrzebuje wzorca czasu jednej sekundy. Posłużymy się tutaj dobrze znanym i używanym chyba w każdym odcinku modulem **StrobeGenerator**. Za pomocą parametru **PERIOD_US** skonfigurujemy go tak, aby dokładnie co jedną sekundę ustawiał swoje wyjście **Strobe_o** w stan wysoki na jeden takt zegarowy.

Wyjście tego modułu pełni dwie funkcje. Pierwszą z nich jest sterowanie multiplexerem **Counter_25_I_0** (automat nazwał multiplexer licznikiem...?!), którego wejście **d0** połączone jest z wyjściem multiplexera **mux_6**, omawianego wcześniej, a wejście **d1** połączone jest z masą, czyli z liczbą zero. Wyjście tego multiplexera doprowadzone jest do wejścia licznika **Counter**, który de facto stanowi zbiór 26 przerzutników D połączonych równolegle. Te przerzutniki, z każdym zboczem rosnącym sygnału zegarowego, wpisują do swojej pamięci: wartość licznika powiększoną o 1, aktualną wartość licznika (czyli nie zmieniają swojego stanu) lub same zero. Ostatnia wymieniona możliwość ma zastosowanie



Rysunek 1. Schemat układu wygenerowany automatycznie na podstawie kodu w języku Verilog

```
// Plik frequency_meter.v
`default_nettype none

module FrequencyMeter #(
    parameter CLOCK_HZ = 25_000_000
)(
    input wire Clock,           // Pin 20
    input wire Reset,           // Pin 17
    input wire SignalAsync_i,   // Pin 75
    output wire [7:0] Cathodes_o, // Pin 40 41 42 43 45 47 51 52
    output wire [7:0] Segments_o // Pin 39 38 37 36 35 34 30 29
);

// Synchronizacja sygnału wejściowego z domeną zegarową
wire SignalSync;

Synchronizer #(
    .WIDTH(1)
) Synchronizer_inst(
    .Clock(Clock),
    .Reset(Reset),
    .Async_i(SignalAsync_i),
    .Sync_o(SignalSync)
);

// Rozpoznawanie zbocza rosnącego badanego sygnału
wire SignalEdge;

EdgeDetector EdgeDetector_inst(
    .Clock(Clock),
    .Reset(Reset),
    .Signal_i(SignalSync),
    .RisingEdge_o(SignalEdge),
    .FallingEdge_o()
);

// Generowanie okna pomiarowego o długości 1 sekundy
wire OneSecondStrobe;

StrobeGenerator #(
    .CLOCK_HZ(CLOCK_HZ),
    .PERIOD_US(1_000_000)
) StrobeGenerator_inst(
    .Clock(Clock),
    .Reset(Reset),
    .Enable_i(1'b1),
    .Strobe_o(OneSecondStrobe)
);

// Licznik do zliczania zboczy rosnących mierzonego sygnału
// w oknie pomiarowym o długości 1 sekundy
reg [25:0] Counter; // 1

always @(posedge Clock, negedge Reset) begin
    if(!Reset)
        Counter <= 0;
    else if(OneSecondStrobe)
        Counter <= 0;
    else if(SignalEdge)
        Counter <= Counter + 1'b1;
end

// Konwersja wyniku na format BCD
wire [31:0] Decimal; // 2

DoubleDabble #(
    .INPUT_BITS(26),
    .OUTPUT_DIGITS(8) // 3
) DoubleDabble_inst(
    .Clock(Clock),
    .Reset(Reset),
    .Start_i(OneSecondStrobe),
    .Busy_o(),
    .Done_o(),
    .Binary_i(Counter),
    .BCD_o(Decimal) // 4
);

// Instancja sterownika wyświetlacza
DisplayMultiplex #(
    .CLOCK_HZ(CLOCK_HZ),
    .SWITCH_PERIOD_US(1000),
    .DIGITS(8)
) DisplayMultiplex_inst(
    .Clock(Clock),
    .Reset(Reset),
    .Data_i(Decimal),
    .DecimalPoints_i(8'b00000000),
    .Cathodes_o(Cathodes_o),
    .Segments_o(Segments_o),
    .SwitchCathode_o()
);

endmodule

`default_nettype wire
```

Listing 1. Kod pliku frequency_meter.v

wtedy, kiedy kończy się czas pomiaru i cała operacja zaczyna się od nowa. Wróćmy do multipleksa i licznika: kiedy wyjście z modułu **StrobeGenerator** znajduje się w stanie wysokim, to w kolejnym takcie zegarowym licznik **Counter** zostanie zresetowany, ponieważ multipleks **Counter_25_I_O** na wejście licznika poda same zera.

Drugie zastosowanie sygnału generowanego przez moduł **StrobeGenerator** obejmuje uruchamianie konwersji liczby binarnej z wyjścia licznika na format dziesiętny BCD, aby pokazać

ją na wyświetlaczu LED w formacie zrozumiałym dla człowieka. Sygnał z wyjścia **StrobeGenerator** jest zatem doprowadzony do wejścia **Start_i** modułu **DoubleDabble** w wersji sekwencyjnej (poznaliśmy go w 22 odcinku kursu, opublikowanym w EP 08/2024).

Może trochę dziwić, że jeden sygnał powoduje jednocześnie zerowanie licznika i konwersję jego zawartości na format dziesiętny. Jednak nie ma w tym nic nieprawidłowego! Ustawienie stanu wysokiego na wyjściu modułu **StrobeGenerator** uruchamia dwie operacje (które wykonywane są jednocześnie w tym samym takcie zegarowym):

1. Zerowanie licznika, tzn. po wystąpieniu zbocza rosnącego sygnału zegarowego, stan tego licznika zostanie ustawiony na zero.
2. Uruchomienie konwersji stanu licznika na format BCD, tzn. po wystąpieniu zbocza rosnącego sygnału zegarowego, obecny stan licznika zostanie odczytany przez moduł **DoubleDabble** i skopiowany do jego wewnętrznego rejestru, a następnie rozpocznie się konwersja.

Dochodzimy do ostatniego modułu, czyli sterownika wyświetlacza **DisplayMultiplex**. Do jego wejścia **Data_i** doprowadzono wprost wynik z wyjścia **BCD_o** modułu **DoubleDabble**. Moduł ten nie potrzebuje żadnych sygnałów wyzwalających, więc zmiana na wejściu danych znajduje od razu odzwierciedlenie na wyjściach sterujących katodami i segmentami wyświetlacza multipleksowanego LED.

Moduł FrequencyMeter

Kod modułu **FrequencyMeter** pokazano na **listingu 1**. Synteza tego modułu daje w rezultacie schemat, którego działanie omówiliśmy powyżej.

Jedyna kwestia, jaka nie została dotychczas wyjaśniona, to powód, dla którego licznik **Counter** ma długość 26 bitów (linia 1). Wartość ta wynika z faktu, że wyświetlacz zamontowany na płytce User Interface Board jest 8-cyfrowy. Maksymalna liczba, którą możemy zapisać w tego typu liczniku, to $2^{26}-1$, czyli 67108863, zatem taka właśnie liczba zmieści się na wyświetlaczu. Gdyby licznik był 27-bitowy, wówczas potrzebowalibyśmy wyświetlacza 9-cyfrowego. Teoretycznie 67,108863 MHz jest największą częstotliwością, jaką można zmierzyć opisywanym miernikiem. W rzeczywistości jest ona dużo mniejsza (około 10 MHz), bo ogranicza nas zegar w FPGA o częstotliwości 25 MHz oraz pojemność pasywności ścieżek na płytce, które nie są przystosowane do przenoszenia sygnałów o wymienionej wcześniej częstotliwości.

Rzućmy okiem na konfigurację parametrów modułu **DoubleDabble**. W linii 3 informujemy moduł, że zmienna wejściowa ma 26 bitów. Następnie w linii 4 ustalamy, że na wyjściu chcemy mieć 8 cyfr BCD, a każda z nich ma 4 bity, więc do odczytania wyniku potrzebujemy 32-bitowej zmiennej typu wire (linia 2).

Testbench modułu FrequencyMeter

Omówimy teraz kod testera (**listing 2**), aby sprawdzić, czy miernik częstotliwości działa prawidłowo. Metodologia testbenchu będzie niezwykle prosta. Utworzymy generator sygnału zegarowego, który zostanie doprowadzony na wejście miernika, a następnie zatrzymamy sekwencję testową, czekając, aż zostanie podany wynik pomiaru.

Na początku kodu przeprowadzamy konfigurację częstotliwości sygnału zegarowego za pomocą parametru **CLOCK_HZ**, tak jak to robiliśmy w poprzednich odcinkach. Poniżej, w linii 1, ustalamy częstotliwość badanego sygnału za pomocą parametru **TEST_SIGNAL_HZ**.

Generatory sygnału zegarowego i testowego zrealizowane są bardzo podobnie. Musimy utworzyć zmienną **TestSignal** typu reg (linia 2), inicjalizując ją zerem lub jedynką, a następnie w bloku


```
// Plik frequency_meter_tb.v
`timescale 1ns/1ns
`default_nettype none
module FrequencyMeter_tb();

    // Konfiguracja
    parameter CLOCK_HZ = 2_000_000;
    parameter TEST_SIGNAL_HZ = 123_456; // 1

    // Generator sygnału zegarowego
    reg Clock = 1'b1;
    always begin
        #((1_000_000_000.0 / (2 * CLOCK_HZ)));
        Clock = !Clock;
    end

    // Generator testowego sygnału do pomiaru
    reg TestSignal = 1'b0; // 2
    always begin // 3
        #((1_000_000_000.0 / (2 * TEST_SIGNAL_HZ)));
        TestSignal = !TestSignal;
    end

    // Eksport wyników symulacji do pliku
    initial begin
        $dumpfile("frequency_meter.vcd");
        $dumpvars(0, FrequencyMeter_tb);
    end

    // Instancja testowanego modułu
    FrequencyMeter #(
        .CLOCK_HZ(CLOCK_HZ)
    ) DUT(
        .Clock(Clock),
        .Reset(Reset),
        .SignalAsync_i(TestSignal),
        .Cathodes_o(), // 4
        .Segments_o(), // 5
        // 6
    );

    reg Reset = 0;

    // Sekwencja testowa
    initial begin
        $timeformat(-9, 3, "ns", 10);
        $display("==== START ====");
        $display("Test freq: %0d", TEST_SIGNAL_HZ); // 7

        // Uruchomienie modułu
        @(posedge Clock);
        Reset = 1'b1; // 8

        // Oczekiwanie na wynik
        @(posedge DUT.DoubleDabble_inst.Done_o) // 9

        repeat(100) @(posedge Clock); // 10

        $display("Result: %0H", DUT.Decimal); // 11

        $display("==== END ====");
        $finish;
    end

endmodule
```

Listing 2. Kod pliku frequency_meter_tb.v

always (linia 3) odwracamy stan tej zmiennej co pewien czas, zależny od parametru **TEST_SIGNAL_HZ**.

Przejdźmy do instancji testowanego modułu, który zwyczajowo nazywamy **DUT** (*device under test*). Do jego wejścia **SignalAsync_i** podłączamy testowy sygnał **TestSignal** (linia 4). Zwróć uwagę, że wyjścia **Cathodes_o** oraz **Segments_o** nie zostały podłączone do niczego. Taka decyzja powodowana jest faktem, że zmiany sygnałów sterujących poszczególnymi segmentami wyświetlacza są trudne do przeanalizowania dla człowieka, więc nie

```
@echo off
iverilog -o frequency_meter.o ^
frequency_meter.v ^
frequency_meter_tb.v ^
edge_detector.v ^
double_dabble.v ^
synchronizer.v ^
display_multiplex.v ^
decoder_7seg.v ^
strobe_generator.v ^
vvp frequency_meter.o
del frequency_meter.o
```

Listing 3. Kod pliku frequency_meter.bat

```
VCD info: dumpfile frequency_meter.vcd opened for output.
==== START ====
Test freq: 123456
Result: 123457
==== END =====
frequency_meter_tb.v:62: $finish called at 1000076500 (1ns)
```

Listing 4. Log symulacji na konsoli

będziemy się nimi interesować. Moduł wyświetlacza testowaliśmy już w 9 odcinku kursu i stąd wiemy, że działa. Wystarczy tylko, że będziemy obserwować wynik bezpośrednio z wyjścia miernika częstotliwości.

Zobaczmy sekwencję testową. W linii 7 printujemy na konsoli informację o częstotliwości testowego sygnału. Ta operacja posłuży do weryfikacji, czy moduł daje poprawny wynik. W tym odcinku nie będziemy przeprowadzać automatycznej weryfikacji w testbenchu, choć nie byłoby to trudne – spróbuj sam dodać taką funkcjonalność.

W linii 8 ustawiamy zmienną **Reset** w stan wysoki, co powoduje rozpoczęcie pracy modułu **FrequencyCounter** oraz wszystkich jego modułów podrzędnych. Moduł miernika częstotliwości pracuje, zliczając wszystkie zbocza testowanego sygnału. Musimy poczekać, aż skończy się okno pomiarowe i moduł zwróci wynik pomiaru. W linii 9 oczekujemy na zbocze rosnące wyjścia **Done_o** w module **DoubleDabble_inst**, który jest wewnętrznym modułem w **DUT** (linia 9). Za pomocą operatora kropki możemy sięgać do zmiennych leżących wewnątrz zagnieżdżonych modułów.

W linii 9 czekamy 100 cykli zegarowych (linia 10) – celem narysowania odstępu w przeglądarce GTKWave, aby lepiej było widać wynik. Na końcu, w linii 11, printujemy na konsoli wynik pomiaru. Zwróć uwagę na drobną różnicę między liniami 7 i 11. W linii 7 odczytujemy parametr, który zapisany jest w pamięci w postaci binarnej i chcemy go wyświetlić w formacie dziesiętnym – dlatego stosujemy operator **%0d** (działa on podobnie jak składnia funkcji **printf()** z C++). Natomiast w linii 11 odczytujemy wynik w formacie BCD i – chcąc wyświetlić go na konsoli – musimy posłużyć się operatorem **%0H**, który zazwyczaj służy do wyświetlania liczb w postaci szesnastkowej.

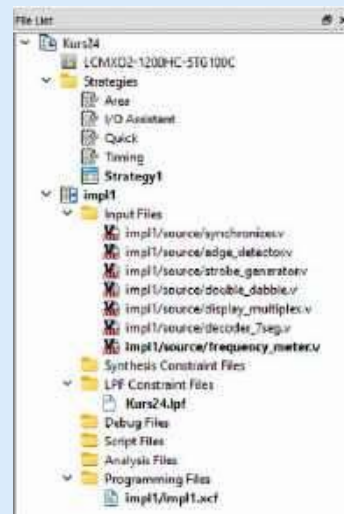
Kod skryptu do uruchomienia symulacji pokazano na **listingu 3**. Zapisy, które zobaczymy na konsoli po zakończeniu symulacji, zaprezentowano na **listingu 4**.

Częstotliwość testowego sygnału i wynik pomiaru różnią się o jeden! Czy moduł działa nieprawidłowo? Nie do końca. Metoda pomiaru, którą stosujemy, z założenia dopuszcza możliwość różnicy o jeden. Dzieje się tak dlatego, że w oknie pomiarowym występuje jakaś całkowita liczba zboczy badanego sygnału. Może się zdarzyć (i bardzo często tak jest), że okno pomiarowe kończy się w trakcie okresu badanego sygnału. Zatem istnieje możliwość, że w kolejnych oknach pomiarowych policzymy jedno więcej lub jedno mniej zboczy badanego sygnału.

Testujemy na żywo

Czas na ćwiczenia praktyczne w prawdziwym układzie FPGA. Uruchom program Lattice Diamond, utwórz nowy projekt i dodaj do niego pliki, które pokazano na **rysunku 2**. Wszystkie pliki możesz ściągnąć z repozytorium na GitHubie <https://github.com/leonow32/verilog-fpga>, a cały gotowy projekt w Diamond pobierzesz z adresu <https://tiny.pl/9pkw6h-v>.

A gdzie jest plik top.v? Nie ma! Tym razem stwierdziłem, że nie ma sensu go tworzyć, ponieważ znalazłaby się w nim wyłącznie instancja modułu **FrequencyMeter** i nic więcej. W takiej sytuacji Diamond rozpoznaje moduł nadrzędny po tym,



Rysunek 2. Drzewko projektu

że w żadnym innym pliku nie ma jego instancji. Nazwa pliku, który został rozpoznany jako nadrzędny, jest pogrubiona w drzewku projektowym.

Przeprowadź syntezę, a następnie otwórz Spreadsheet i skonfiguruj piny układu FPGA, tak jak to pokazano na **rysunku 3**.

W dzisiejszym odcinku nie używamy generatora RC typu OSCH, wbudowanego w FPGA, lecz używamy zewnętrznego generatora kwarcowego, ponieważ jest on dużo bardziej precyzyjny. W takiej sytuacji musimy także określić, jaka jest częstotliwość zegara. Klikamy zakładkę Timing Preferences na dole okna Spreadsheet, po czym ustawiamy wszystko tak, jak to widać na **rysunku 4**.

Generujemy bitstream i wgrywamy do FPGA. Na wyświetlaczu cyfrowym powinna pojawić się liczba 0, jeżeli żaden sygnał nie jest podłączony. Sygnał z dowolnego generatora należy podłączyć do wejścia Rx na złączu goldpin J5, a masę generatora należy podłączyć do wejścia GND. Pamiętaj, że napięcie mierzonego sygnału powinno się zmieniać w zakresie do 0 V do 3,3 V. Film prezentujący działanie układu jest dostępny pod adresem <https://youtu.be/FtbOKaUsGUE>.

W następnym odcinku poznamy, jak za pomocą FPGA zbudować przetwornik cyfrowo-analogowy. Na kolejnym etapie użyjemy pamięci EBR, by cyklicznie odczytywać z niej wartości napięcia, które następnie będziemy przekazywać do przetwornika cyfrowo-analogowego. W taki sposób można wygenerować sygnał analogowy o zupełnie dowolnym kształcie. Poznamy także, w jaki sposób stworzyć zmienny dzielnik częstotliwości, aby móc regulować częstotliwość generowanego sygnału.

Dominik Bieczyński
leonow32@gmail.com

Spreadsheet View															
	Name	Source	Pin	IOB	IOB_VCC	IOB_VDD	IOB_TYPE	PULLDOWN	DRIVE	SENSE	CLOCK	OPEN Drain	SUPPRESSOR	DIRECTION	HYSTERIS
1	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
2	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
3	2	2	2	2	2	2	2	2	2	2	2	2	2	2	2
4	3	3	3	3	3	3	3	3	3	3	3	3	3	3	3
5	4	4	4	4	4	4	4	4	4	4	4	4	4	4	4
6	5	5	5	5	5	5	5	5	5	5	5	5	5	5	5
7	6	6	6	6	6	6	6	6	6	6	6	6	6	6	6
8	7	7	7	7	7	7	7	7	7	7	7	7	7	7	7
9	8	8	8	8	8	8	8	8	8	8	8	8	8	8	8
10	9	9	9	9	9	9	9	9	9	9	9	9	9	9	9
11	10	10	10	10	10	10	10	10	10	10	10	10	10	10	10
12	11	11	11	11	11	11	11	11	11	11	11	11	11	11	11
13	12	12	12	12	12	12	12	12	12	12	12	12	12	12	12
14	13	13	13	13	13	13	13	13	13	13	13	13	13	13	13
15	14	14	14	14	14	14	14	14	14	14	14	14	14	14	14
16	15	15	15	15	15	15	15	15	15	15	15	15	15	15	15
17	16	16	16	16	16	16	16	16	16	16	16	16	16	16	16
18	17	17	17	17	17	17	17	17	17	17	17	17	17	17	17
19	18	18	18	18	18	18	18	18	18	18	18	18	18	18	18
20	19	19	19	19	19	19	19	19	19	19	19	19	19	19	19
21	20	20	20	20	20	20	20	20	20	20	20	20	20	20	20
22	21	21	21	21	21	21	21	21	21	21	21	21	21	21	21
23	22	22	22	22	22	22	22	22	22	22	22	22	22	22	22
24	23	23	23	23	23	23	23	23	23	23	23	23	23	23	23
25	24	24	24	24	24	24	24	24	24	24	24	24	24	24	24
26	25	25	25	25	25	25	25	25	25	25	25	25	25	25	25
27	26	26	26	26	26	26	26	26	26	26	26	26	26	26	26
28	27	27	27	27	27	27	27	27	27	27	27	27	27	27	27
29	28	28	28	28	28	28	28	28	28	28	28	28	28	28	28
30	29	29	29	29	29	29	29	29	29	29	29	29	29	29	29
31	30	30	30	30	30	30	30	30	30	30	30	30	30	30	30
32	31	31	31	31	31	31	31	31	31	31	31	31	31	31	31
33	32	32	32	32	32	32	32	32	32	32	32	32	32	32	32
34	33	33	33	33	33	33	33	33	33	33	33	33	33	33	33
35	34	34	34	34	34	34	34	34	34	34	34	34	34	34	34
36	35	35	35	35	35	35	35	35	35	35	35	35	35	35	35
37	36	36	36	36	36	36	36	36	36	36	36	36	36	36	36
38	37	37	37	37	37	37	37	37	37	37	37	37	37	37	37
39	38	38	38	38	38	38	38	38	38	38	38	38	38	38	38
40	39	39	39	39	39	39	39	39	39	39	39	39	39	39	39
41	40	40	40	40	40	40	40	40	40	40	40	40	40	40	40
42	41	41	41	41	41	41	41	41	41	41	41	41	41	41	41
43	42	42	42	42	42	42	42	42	42	42	42	42	42	42	42
44	43	43	43	43	43	43	43	43	43	43	43	43	43	43	43
45	44	44	44	44	44	44	44	44	44	44	44	44	44	44	44
46	45	45	45	45	45	45	45	45	45	45	45	45	45	45	45
47	46	46	46	46	46	46	46	46	46	46	46	46	46	46	46
48	47	47	47	47	47	47	47	47	47	47	47	47	47	47	47
49	48	48	48	48	48	48	48	48	48	48	48	48	48	48	48
50	49	49	49	49	49	49	49	49	49	49	49	49	49	49	49
51	50	50	50	50	50	50	50	50	50	50	50	50	50	50	50
52	51	51	51	51	51	51	51	51	51	51	51	51	51	51	51
53	52	52	52	52	52	52	52	52	52	52	52	52	52	52	52
54	53	53	53	53	53	53	53	53	53	53	53	53	53	53	53
55	54	54	54	54	54	54	54	54	54	54	54	54	54	54	54
56	55	55	55	55	55	55	55	55	55	55	55	55	55	55	55
57	56	56	56	56	56	56	56	56	56	56	56	56	56	56	56
58	57	57	57	57	57	57	57	57	57	57	57	57	57	57	57
59	58	58	58	58	58	58	58	58	58	58	58	58	58	58	58
60	59	59	59	59	59	59	59	59	59	59	59	59	59	59	59
61	60	60	60	60	60	60	60	60	60	60	60	60	60	60	60
62	61	61	61	61	61	61	61	61	61	61	61	61	61	61	61
63	62	62	62	62	62	62	62	62	62	62	62	62	62	62	62
64	63	63	63	63	63	63	63	63	63	63	63	63	63	63	63
65	64	64	64	64	64	64	64	64	64	64	64	64	64	64	64
66	65	65	65	65	65	65	65	65	65	65	65	65	65	65	65
67	66	66	66	66	66	66	66	66	66	66	66	66	66	66	66
68	67	67	67	67	67	67	67	67	67	67	67	67	67	67	67
69	68	68	68	68	68	68	68	68	68	68	68	68	68	68	68
70	69	69	69	69	69	69	69	69	69	69	69	69	69	69	69
71	70	70	70	70	70	70	70	70	70	70	70	70	70	70	70
72	71	71	71	71	71	71	71	71	71	71	71	71	71	71	71
73	72	72	72	72	72	72	72	72	72	72	72	72	72	72	72
74	73	73	73	73	73	73	73	73	73	73	73	73	73	73	73
75	74	74	74	74	74	74	74	74	74	74	74	74	74	74	74
76	75	75	75	75	75	75	75	75	75	75	75	75	75	75	75
77	76	76	76	76	76	76	76	76	76	76	76	76	76	76	76
78	77	77	77	77	77	77	77	77	77	77	77	77	77	77	77
79	78	78	78	78	78	78	78	78	78	78	78	78	78	78	78
80	79	79	79	79	79	79	79	79	79	79	79	79	79	79	79
81	80	80	80	80	80	80	80	80	80	80	80	80	80	80	80
82	81	81	81	81	81	81	81	81	81	81	81	81	81	81	81
83	82	82	82	82	82	82	82	82	82	82	82	82	82	82	82
84	83	83	83	83	83	83	83	83	83	83	83	83	83	83	83
85	84	84	84	84	84	84	84	84	84	84	84	84	84	84	84
86	85	85	85	85	85	85	85	85	85	85	85	85	85	85	85
87	86	86	86	86	86	86	86	86	86	86	86	86	86	86	86
88	87	87	87	87	87	87	87	87	87	87	87	87	87	87	

Kurs Nordic nRF z BT (4)

Bluetooth LE

W poprzedniej części kursu zajęliśmy się sterowaniem LED-ami naszej płytki deweloperskiej. Dziś dodamy możliwość ustawiania przez smartfon parametrów ich świecenia. Zaprezentujemy również trochę teorii dotyczącej Bluetooth Low Energy.

Bluetooth LE

Na płytce nRF5340DK mamy zamontowany SoC nRF5340, który zawiera dwa procesory Arm® Cortex®-M33: jeden aplikacyjny i jeden sieciowy. Do tej pory działaliśmy na procesorze aplikacyjnym, a teraz zaprzęgniemy procesor sieciowy do obsługi stosu protokołu Bluetooth Low Energy (Bluetooth LE, kolokwialnie BLE).

Bluetooth LE różni się zasadniczo od klasycznego Bluetooth. Powstał w odpowiedzi na rosnące zapotrzebowanie na technologie bezprzewodowe o niskim zużyciu energii, szczególnie w aplikacjach IoT, urządzeniach noszonych (wearables), czujnikach i innych urządzeniach o małej mocy. Tradycyjny Bluetooth (BR/EDR) nie był zoptymalizowany pod kątem energooszczędności, co ograniczało jego zastosowanie w takich aplikacjach. BLE ma, co prawda,

```
#include "led_control.h"
#include <zephyr/kernel.h>
#include <zephyr/drivers/gpio.h>
#include <zephyr/logging/log.h>

LOG_MODULE_REGISTER(main, LOG_LEVEL_DBG);

static const struct gpio_dt_spec buttons[] = {
    GPIO_DT_SPEC_GET(DT_NODELABEL(button0), gpios),
    GPIO_DT_SPEC_GET(DT_NODELABEL(button1), gpios),
    GPIO_DT_SPEC_GET(DT_NODELABEL(button2), gpios),
    GPIO_DT_SPEC_GET(DT_NODELABEL(button3), gpios)
};

static struct gpio_callback cb_data;

static void button_callback(const struct device *dev,
    struct gpio_callback *cb, uint32_t pins)
{
    LOG_INF("Pin mask 0x%08x", pins);

    if (pins & BIT(buttons[0].pin)) {
        led_send(&led_msg){ LED_STATE, 0 };
    } else if (pins & BIT(buttons[1].pin)) {
        led_send(&led_msg){ LED_STATE, 1 };
    } else if (pins & BIT(buttons[2].pin)) {
        led_send(&led_msg){ LED_BLINK, 500 };
    } else if (pins & BIT(buttons[3].pin)) {
        led_send(&led_msg){ LED_BLINK, 100 };
    }
}

int main(void) {
    gpio_init_callback(&cb_data, button_callback,
        BIT(buttons[0].pin) | BIT(buttons[1].pin) |
        BIT(buttons[2].pin) | BIT(buttons[3].pin));

    for (int i = 0; i < ARRAY_SIZE(buttons); ++i) {
        gpio_pin_configure_dt(&buttons[i], GPIO_INPUT);
        gpio_pin_interrupt_configure_dt(&buttons[i],
            GPIO_INT_EDGE_FALLING);
        gpio_add_callback(buttons[i].port, &cb_data);
    }

    int counter = 0;
    while (1) {
        LOG_INF("Tick %d", counter++);
        k_msleep(1000);
    }
    return 0;
}
```

Listing 1. Plik src/main.c

```
#pragma once

typedef struct {
    enum {
        LED_STATE, // param: 0 (off) or 1 (on)
        LED_BLINK // param: time [ms] for every toggle
    } command;
    unsigned int param;
} led_msg;

void led_send(led_msg msg);
```

Listing 2. Plik src/led_control.h



Poprzednie odcinki znajdują się pod adresem:
<https://ulubionykiosk.pl/media>

mniejszy transfer danych (do 1,4 Mbps), ale oferuje dłuższy czas pracy na baterii.

Ponadto Bluetooth LE – poza połączeniami Point-to-Point – potrafi pracować w topologiach Broadcast i Mesh. Pozwala to budować urządzenia, które po prostu rozgłaszają dane, a nawet budować całe sieci urządzeń oparte na BLE.

Tradycyjny Bluetooth i BLE nie są ze sobą kompatybilne. Ponieważ nRF5340 wspiera jedynie Bluetooth LE, do sterowania naszą płytką jest wymagany smartfon lub tablet wspierający tę technologię. Większość urządzeń wyprodukowanych po 2011 roku powinno to wsparcie oferować.

Przedstawimy teraz nieco podstawowych pojęć związanych z BLE. Nie są one kluczowe dla naszego celu, ale powinny znacząco ułatwić zgłębianie tematu i czytanie dokumentacji.

Urządzenia BLE pracują z reguły w jednej z dwóch ról: *Central* lub *Peripheral*.

Central skanuje i inicjuje połączenia z urządzeniami typu *Peripheral*. Pełni funkcję hosta, odpowiada za typowe zadania tej roli, takie jak nawiązywanie połączeń i zarządzanie nimi oraz przetwarzanie danych. W naszym przypadku smartfon będzie pracował w tym trybie.

Peripheral to urządzenie, które rozgłasza swoje dane, umożliwia nawiązanie połączenia oraz udostępnia jeden lub więcej serwisów (*Service*). Pełni funkcję serwera. Urządzenia IoT o ograniczonych zasobach, które wymagają niskiego zużycia energii, zazwyczaj odgrywają właśnie tę rolę – w tym kursie funkcję *Peripheral* będzie pełniła nasza płytka deweloperska.

W przypadku pracy rozgłoszeniowej (*Broadcast oriented*) rolę są podobne: zamiast *Peripheral* mamy *Broadcaster*, a zamiast *Central* jest *Observer*, z tą różnicą, że pierwszy nie udostępnia, a drugi nie nawiązuje połączenia.

Jedno urządzenie może jednocześnie pełnić więcej niż jedną funkcję: udostępniać dane więcej niż jednemu urządzeniu *Central* i odbierać dane z wielu urządzeń typu *Peripheral*.

Wymiana danych w BLE odbywa się w ramach Serwisów i Charakterystyk.

Charakterystyka (*Characteristic*) to podstawowy fragment informacji/danych, który serwer (*Peripheral*) udostępnia klientowi (*Central*). Charakterystyka jest zawsze częścią serwisu. W naszym przypadku będzie to parametr pracy LED-a.

Serwis (*Service*) to po prostu grupa charakterystyk. Z reguły łączy się one w logiczną całość – np. temperatura, wilgotność i jakość powietrza. My utworzymy jeden serwis związany ze sterowaniem LED-em.

W poprzednim odcinku...

...stworzyliśmy aplikację, w której mogliśmy włączać, wyłączać oraz migać diodą naszej płytki deweloperskiej nRF5340-DK za pomocą przycisków.

Projekt składał się z trzech plików źródłowych (listyngi 1...3). Dodatkowo mieliśmy plik CMakeLists.txt do budowania

projektu (listing 4), oraz konfigurację projektu w pliku `prj.conf` (listing 5). Wszystkie pliki źródłowe wymienione w tym odcinku kursu można znaleźć w materiałach dodatkowych do artykułu na stronie <https://ep.com.pl>.

Włączamy Bluetooth

Ustaliliśmy już, że nasza płytką będzie pracowała w standardzie BLE jako *Peripheral*. Należy zatem włączyć taką możliwość w `prj.conf` – pliku konfiguracyjnym projektu – poprzez dodanie trzech linii, widocznych na listingu 6.

Włączenie i rozpoczęcie rozgłaszania BLE obsługuje nowy plik `src/bt_control.c` (listing 7).

Plik nagłówkowy (`src/bt_control.h`) będzie bardzo prosty (listing 8).

Funkcja `bt_control_init()` powinna być wywołana w funkcji `main()`, po uprzednim inkludowaniu pliku nagłówkowego `bt_control.h`. Ma za zadanie jedynie wywołać Zephyrową funkcję `bt_enable()` – która rozpocznie proces włączania funkcjonalności Bluetooth. Gdy proces się zakończy, automatycznie wywołana zostanie `bt_ready_cb()`.

Funkcja `bt_ready_cb()` ma z kolei wyświetlić na konsoli adres MAC płytki oraz rozpocząć rozgłaszanie, aby można było odnaleźć płytkę za pomocą skanowania w smartfonie.

Na końcu rejestrowane są funkcje, które wywołane będą w przypadku podłączenia się i odłączenia się smartfona od nRF5340-DK.

Należy jeszcze pamiętać, aby dodać plik `src/bt_led_svc.c` do *build systemu* w pliku `CMakeLists.txt`.

Zbudujemy projekt i wgrajmy kod na płytkę

Aby nieco uprościć pracę, będziemy używać natywnej konfiguracji płytki nRF5340-DK. Wybierając konfigurację budowania przez przycisk *Add build configuration*, należy wybrać `nrf5340dk_nrf5340_cpuapp` lub – jeśli budujecie projekt poprzez konsolę – należy wpisać `west build -b nrf5340dk_nrf5340_cpuapp`. Pliki z poprzedniego odcinka kursu z katalogu *boards* nie będą nam potrzebne.

Jak już wspomnieliśmy, projekt na SoC nRF5340 z BLE używa obu procesorów. Jak zatem oprogramować procesor sieciowy i wgrać na niego firmware?

W niektórych (starszych) kursach lub wątkach w Internecie pojawia się polecenie, aby najpierw skompilować i wgrać *firmware* obsługujący BLE na procesor sieciowy, a potem – swój aplikacyjny *firmware* na procesor aplikacyjny. Nie jest to obecnie wymagane, gdyż narzędzie *west* zbuduje dla Was plik `hex` zawierający firmware obu procesorów i automatycznie wgra oba.

Pod koniec budowania w konsoli można zobaczyć wpis „Generating zephyr/merged_domains.hex”. Tu właśnie jest tworzony plik z *firmware* dla obu procesorów.

```
cmake_minimum_required(VERSION 3.20.0)
find_package(Zephyr REQUIRED HINTS $ENV{ZEPHYR_BASE})

project(thread_test)

target_sources(app PRIVATE
    src/main.c
    src/led_control.c
)
```

Listing 4. Plik `/CMakeLists.txt`

```
#include "led_control.h"
#include <zephyr/kernel.h>
#include <zephyr/drivers/gpio.h>
#include <zephyr/logging/log.h>

LOG_MODULE_REGISTER(led_control, LOG_LEVEL_DBG);

static const struct gpio_dt_spec led =
    GPIO_DT_SPEC_GET(DT_NODELABEL(led0), gpios);

K_MSGQ_DEFINE(led_queue, sizeof(led_msg), 10, 1);

void led_send(led_msg msg) {
    k_msgq_put(&led_queue, &msg, K_NO_WAIT);
}

static void led_worker(void *a, void *b, void *c) {
    gpio_pin_configure_dt(&led, GPIO_OUTPUT_ACTIVE);

    led_msg msg;
    k_timeout_t timeout = K_FOREVER;

    while (1) {
        if (k_msgq_get(&led_queue, &msg, timeout) == 0) {
            if (msg.command == LED_STATE) {
                LOG_INF("LED_STATE %d received", msg.param);
                gpio_pin_set_dt(&led, msg.param);
                timeout = K_FOREVER;
            } else if (msg.command == LED_BLINK) {
                LOG_INF("LED_BLINK %dms received", msg.param);
                timeout = K_MSEC(msg.param);
            }
        } else {
            gpio_pin_toggle_dt(&led);
        }
    }
}

K_THREAD_DEFINE(led_thread, 500, led_worker, NULL, NULL, NULL,
    K_HIGHEST_THREAD_PRIO, 0, 0);

#ifdef CONFIG_SHELL

#include <zephyr/shell/shell.h>
#include <stdlib.h>

static int cmd_led_on(const struct shell *shell, size_t argc, char **argv) {
    led_send(led_msg){ LED_STATE, 1 };
    shell_print(shell, "LED on");
    return 0;
}

static int cmd_led_off(const struct shell *shell, size_t argc, char **argv) {
    led_send(led_msg){ LED_STATE, 0 };
    shell_print(shell, "LED off");
    return 0;
}

static int cmd_led_blink(const struct shell *shell, size_t argc, char **argv) {
    if (argc != 2) {
        shell_print(shell, "led blink <time in ms>");
        return -EINVAL;
    }

    uint32_t period = strtoul(argv[1], NULL, 10);
    led_send(led_msg){ LED_BLINK, period };
    shell_print(shell, "LED blinking period %d ms:", period);
    return 0;
}

SHELL_STATIC_SUBCMD_SET_CREATE(led_menu,
    SHELL_CMD(on, NULL, "LED on", cmd_led_on),
    SHELL_CMD(off, NULL, "LED off", cmd_led_off),
    SHELL_CMD(blink, NULL, "LED blink with period [ms]", cmd_led_blink),
    SHELL_SUBCMD_SET_END
);

SHELL_CMD_REGISTER(led, &led_menu, "LED", NULL);
#endif
```

Listing 3. Plik `src/led_control.c`

```
CONFIG_GPIO=y

CONFIG_SERIAL=y
CONFIG_CONSOLE=y
CONFIG_UART_CONSOLE=y

CONFIG_LOG=y
CONFIG_LOG_PROCESS_THREAD_SLEEP_MS=100

CONFIG_USE_SEGGER_RTT=y
CONFIG_LOG_BACKEND_RTT=y

CONFIG_SHELL=y
CONFIG_SHELL_BACKEND_SERIAL=y

CONFIG_GPIO_SHELL=y
```

Listing 5. Plik `/prj.conf`

```
CONFIG_BT=y
CONFIG_BT_PERIPHERAL=y
CONFIG_BT_DEVICE_NAME="Led control"
```

Listing 6. Włączanie BLE w pliku `prj.conf`

Włączenie opcji `CONFIG_BT=y` w ustawieniach projektu automatycznie dodaje do projektu budowanie drugiego obrazu, przeznaczonego dla procesora sieciowego, zawierającego wszystkie potrzebne warstwy stosu BLE. W katalogu `/build/hci_ipc` odnajdziesz wynik budowania projektu na procesor sieciowy.

Tego typu projekty są całkiem niezłe wspierane przez wtyczkę nRF Connect. Wystarczy, że w sekcji Applications wybierzesz procesor sieciowy i klikniesz na katalogu wynikowym tuż pod nim, a będziesz mógł przeglądać źródła tego projektu (**rysunek 1**).

Po wgraniu projektu przebieg procesu włączania i rozpoczynania rozgłaszania możemy sprawdzić na logach w konsoli (**listing 9**).

Zauważ wpisy w logu, których nie wprowadzaliśmy – te oznaczone jako `bt_hci_core`. Są to informacje ze stosu BLE Zephyra.

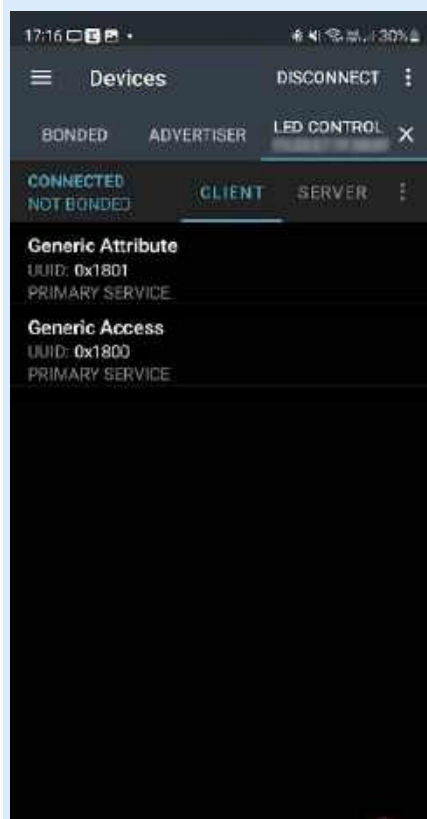
Teraz możesz odnaleźć urządzenie „LED control” w standardowym menu Bluetooth telefonu.

My posłużymy się jednak bardziej zaawansowanym narzędziem.

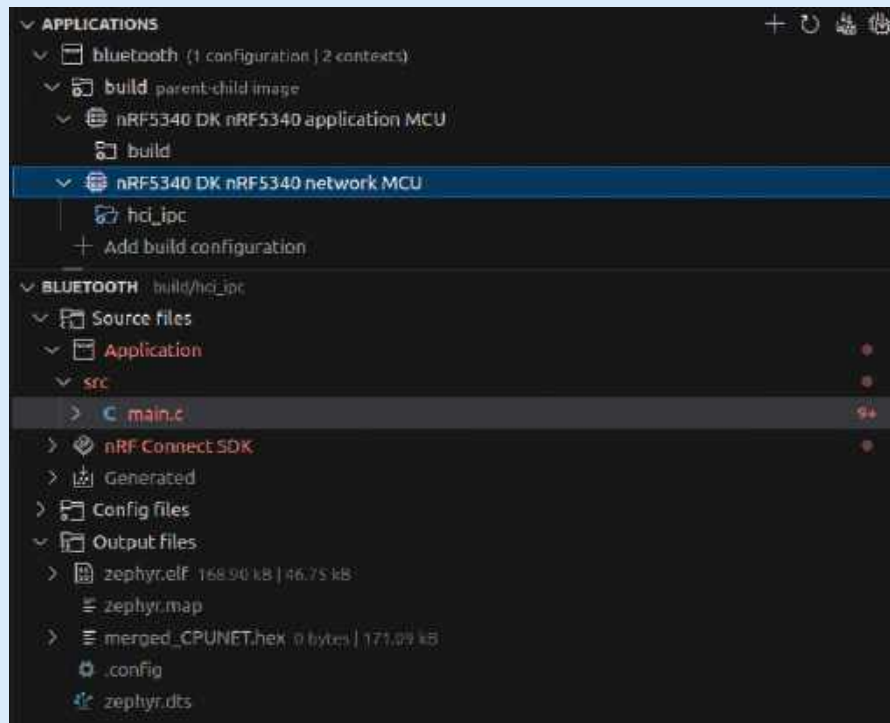
Aplikacja nRF Connect

Zainstaluj „nRF Connect for Mobile” – jest dostępna zarówno na smartfony z Androidem, jak i z iOS.

Po uruchomieniu rozpocznij skanowanie (w prawym górnym rogu). Nasza płytka powinna się pojawić pośród innych urządzeń BT (**rysunek 2**).



Rysunek 2. Ekran aplikacji nRF Connect – wyszukiwanie płytki



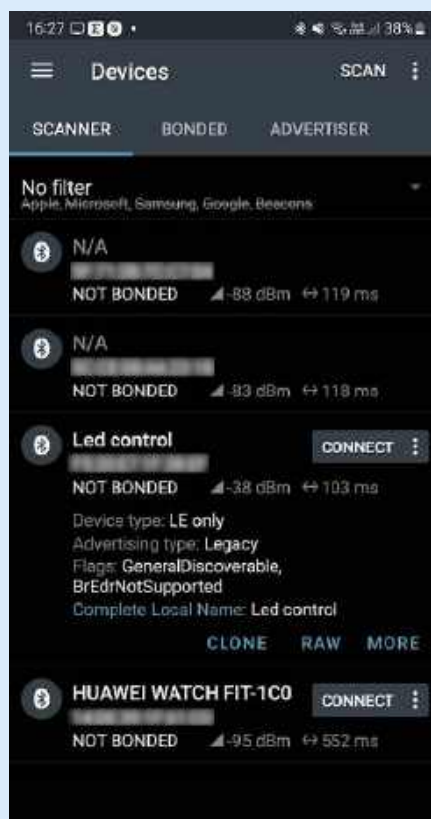
Rysunek 1. Przegląd projektu procesora sieciowego we wtyczce nRFConnect

Po kliknięciu w CONNECT powinniśmy zobaczyć serwisy udostępniane przez nasze urządzenie. Będą to tylko podstawowe dane, potrzebne np. do pobrania nazwy naszego urządzenia (**rysunek 3**).

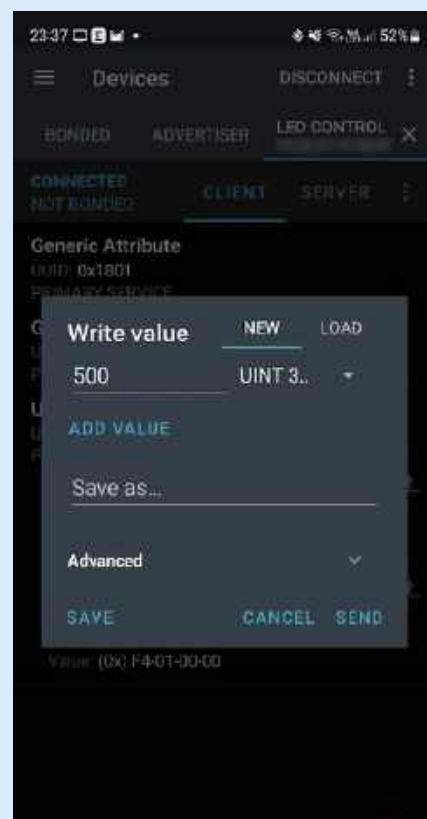
Czas zatem dodać własny serwis.

Serwis LED

Cały serwis obsługujący LED mieści się w pliku `src/bt_led_svc.c` (**listing 10**).



Rysunek 3. Podstawowe serwisy



Rysunek 4. Wpisywanie parametru LED

```

#include <zephyr/bluetooth/bluetooth.h>
#include <zephyr/bluetooth/conn.h>
#include "bt_control.h"

#include <zephyr/logging/log.h>
LOG_MODULE_REGISTER(bt_control, LOG_LEVEL_DBG);

// advertisement data (AdvData)
static const struct bt_data adv_data[] = {
    BT_DATA_BYTES(BT_DATA_FLAGS, (BT_LE_AD_GENERAL | BT_LE_AD_NO_BREDR)),
    BT_DATA(BT_DATA_NAME_COMPLETE, CONFIG_BT_DEVICE_NAME, (sizeof(CONFIG_BT_DEVICE_NAME)-1)),
};

static void connected_cb(struct bt_conn *conn, uint8_t err)
{
    if (err) {
        LOG_ERR("Connection failed. error:%u", err);
        return;
    }

    struct bt_conn_info info;
    int ret = bt_conn_get_info(conn, &info);
    if (ret == 0) {
        char addr_dst[BT_ADDR_LE_STR_LEN];
        bt_addr_le_to_str(info.le.dst, addr_dst, sizeof(addr_dst));
        LOG_INF("Connected to: %s", addr_dst);
    } else {
        LOG_ERR("MAC read failed");
    }
}

static void disconnected_cb(struct bt_conn *conn, uint8_t reason)
{
    LOG_INF("Disconnected. Reason: %u", reason);
}

static struct bt_conn_cb conn_callbacks = {
    .connected = connected_cb,
    .disconnected = disconnected_cb
};

void bt_ready_cb(int status) {

    //MAC
    bt_addr_le_t addrs[CONFIG_BT_ID_MAX];
    size_t count = CONFIG_BT_ID_MAX;
    bt_id_get(addrs, &count);

    //print MAC
    if (count > 0) {
        char addr_str[BT_ADDR_LE_STR_LEN];
        bt_addr_le_to_str(&addrs[0], addr_str, sizeof(addr_str));
        LOG_DBG("BLE MAC: %s", addr_str);
    } else {
        LOG_ERR("Error while getting MAC");
    }

    //Start advertising
    int err = bt_le_adv_start(BT_LE_ADV_CONN, adv_data, ARRAY_SIZE(adv_data), NULL, 0);
    if (err) {
        LOG_ERR("BLE advertisement error: %d", err);
    } else {
        LOG_DBG("BLE advertisement started");
    }

    //Register connected/disconnected callbacks
    bt_conn_cb_register(&conn_callbacks);
}

int bt_control_init(void) {
    LOG_DBG("BLE init...");

    int ret = bt_enable(bt_ready_cb);
    if (ret) {
        LOG_ERR("BLE init error: %d\n", ret);
        return ret;
    }

    return ret;
}

```

Listing 7. Plik src/bt_control.c

Towarzyszący mu plik nagłówkowy src/bt_led_svc.h pozwala na spięcie obu charakterystyk serwisu z wcześniej napisanym kodem (listing 11).

```

#pragma once

int bt_control_init(void);

```

Listing 8. Plik nagłówkowy src/bt_control.h

```

[00:00:03.084,899] <dbg> bt_control: bt_control_init: BLE init...
[00:00:03.086,393] <inf> main: Tick 0
[00:00:03.097,869] <inf> bt_hci_core: HW Platform: Nordic Semiconductor (0x0002)
[00:00:03.097,900] <inf> bt_hci_core: HW Variant: nRF53x (0x0003)
[00:00:03.097,900] <inf> bt_hci_core: Firmware: Standard Bluetooth controller (0x00) Version 54.58864 Build 1214809870
[00:00:03.099,426] <inf> bt_hci_core: Identity: F5:20:E7:1F:38:EF (random)
[00:00:03.099,456] <inf> bt_hci_core: HCI: version 5.4 (0x0d) revision 0x218f, manufacturer 0x0059
[00:00:03.099,487] <inf> bt_hci_core: LMP: version 5.4 (0x0d) subver 0x218f
[00:00:06.844,604] <dbg> bt_control: bt_ready_cb: BLE MAC: F5:20:E7:1F:38:EF (random)
[00:00:06.846,313] <dbg> bt_control: bt_ready_cb: BLE advertisement started

```

Listing 9. Log procesu włączania i rozpoczynanie rozgłaszania

Definicja serwisu odbywa się poprzez makro BT_GATT_SERVICE_DEFINE na końcu pliku.

To makro tworzy serwis z identyfikatorem BT_UUID_LED_SERVICE oraz dwie charakterystyki: jedną do zapisu stanu LED-a i drugą – do zapisu okresu, z jakim ma migać.

UUID (Universally Unique Identifier) to unikalny identyfikator, który w Bluetooth Low Energy służy do rozpoznawania serwisów i charakterystyk. Każdy serwis i charakterystyka

musi mieć przypisany UUID, który umożliwia identyfikację w trakcie komunikacji między urządzeniami. UUID mają długość 128 bitów, jednak w przypadku oficjalnie rozpoznawanych identyfikatorów korzystamy ze skróconych, 16-bitowych odpowiedników. Dla niestandardowych serwisów UUID można wygenerować, np. korzystając z narzędzi online, takich jak [1]. Jednak wygodnie jest, jeśli pierwsza część UUID serwisu kończy się zerem, a charakterystyki mają tę wartość kolejno zwiększaną o jeden.

Konfiguracja naszych dwóch charakterystyk odbywa się poprzez kolejne parametry makra BT_GATT_CHARACTERISTIC:

- UUID: kolejno zwiększane o 1,
- typ: BT_GATT_CHRC_WRITE (charakterystyka do zapisu),
- uprawnienia: BT_GATT_PERM_WRITE (zezwoleń na zapis),
- callback odczytu: brak (więc podano NULL),
- callback zapisu: write_chr_cb,
- dodatkowe dane: brak (więc NULL).

W tym przykładzie funkcje, które będą wywołane przy zapisie do charakterystyki poprzez smartfon (*callback*), to jedna i ta sama funkcja – write_chr_cb, a rozróżnienie następuje w jej środku poprzez porównanie UUID. Po wpisaniu parametru w charakterystyce BT_UUID_LED_STATE_CHR zostanie też uruchomiona funkcja skojarzona z write_state_cb, natomiast jeśli parametr zostanie wpisany do charakterystyki BT_UUID_LED_BLINK_CHR, prześlemy go do write_blink_cb.

Oczywiście należy jeszcze pamiętać, aby dodać plik src/bt_led_svc.c do systemu budowania w pliku CMakeLists.txt.

Firmware w tym stanie będzie już zgłaszał naszą charakterystykę. Będziemy nawet w logu widzieć wysłane ze smartfona dane. Aby faktycznie zmienić stan LED, musimy nieco dodać do pliku src/main.c.

Dodajemy funkcje przekazujące parametr do modułu led_control w pliku main.c (**listing 12**) i przekazujemy te funkcje serwisowi w funkcji main tuż po włączeniu BLE (**listing 13**). Oczywiście trzeba też dołączyć bt_led_svc.h w pliku main.c

Od teraz, korzystając z aplikacji mobilnej, możemy do naszej charakterystyki o UUID 7c8482f2-1104-44b8-a354-3f5b8488083d wpisać nowy okres migania LED w milisekundach (**rysunek 4**), pamiętając o wybraniu typu zmiennej jako UINT 32 (Little Endian). Nowa wartość zostanie natychmiast przesłana do naszej płytki, gdzie możemy w czasie rzeczywistym obserwować jej wpływ na zachowanie diody. Przykładowo, po wpisaniu 500, zobaczymy, że nasz LED co sekundę zaświeca się i gaśnie.

Podsumowanie

W tej części kursu rozszerzyliśmy kod sterujący LED za pomocą przycisków i konsoli o możliwość sterowania diodą LED poprzez

```
#include <zephyr/bluetooth/bluetooth.h>
#include <zephyr/bluetooth/uuid.h>
#include <zephyr/bluetooth/gatt.h>
#include "bt_led_svc.h"

#include <zephyr/logging/log.h>
LOG_MODULE_REGISTER(bt_led_svc, LOG_LEVEL_DBG);

// UUID for LED Service
#define BT_UUID_LED_SERVICE BT_UUID_DECLARE_128 \
    (BT_UUID_128_ENCODE(0x7c8482f0, 0x1104, 0x44b8, 0xa354, 0x3f5b8488083d))
// UUID for LED Characteristics
#define BT_UUID_LED_STATE_CHR BT_UUID_DECLARE_128 \
    (BT_UUID_128_ENCODE(0x7c8482f1, 0x1104, 0x44b8, 0xa354, 0x3f5b8488083d))
#define BT_UUID_LED_BLINK_CHR BT_UUID_DECLARE_128 \
    (BT_UUID_128_ENCODE(0x7c8482f2, 0x1104, 0x44b8, 0xa354, 0x3f5b8488083d))

// Callback function pointers for BLE write operations
static led_cb_t write_state_cb = NULL;
static led_cb_t write_blink_cb = NULL;

static ssize_t write_chr_cb(struct bt_conn *conn,
                           const struct bt_gatt_attr *attr,
                           const void *buf, uint16_t len,
                           uint16_t offset, uint8_t flags) {

    LOG_HEXDUMP_INF(buf, len, "Received buffer:");

    if (len != sizeof(uint32_t)) {
        return BT_GATT_ERR(BT_ATT_ERR_INVALID_ATTRIBUTE_LEN);
    }

    unsigned int received = *(unsigned int*)buf;

    if (bt_uuid_cmp(attr->uuid, BT_UUID_LED_STATE_CHR) == 0) {
        LOG_INF("Led state update: %u", received);
        if (write_state_cb != NULL) {
            write_state_cb(received);
        }
    }

    if (bt_uuid_cmp(attr->uuid, BT_UUID_LED_BLINK_CHR) == 0) {
        LOG_INF("Led blink update: %u", received);
        if (write_blink_cb != NULL) {
            write_blink_cb(received);
        }
    }

    return len;
}

// Define the GATT service and characteristics
BT_GATT_SERVICE_DEFINE(led_service,
    BT_GATT_PRIMARY_SERVICE(BT_UUID_LED_SERVICE),

    BT_GATT_CHARACTERISTIC(BT_UUID_LED_STATE_CHR,
        BT_GATT_CHRC_WRITE, BT_GATT_PERM_WRITE,
        NULL, write_chr_cb, NULL),

    BT_GATT_CHARACTERISTIC(BT_UUID_LED_BLINK_CHR,
        BT_GATT_CHRC_WRITE, BT_GATT_PERM_WRITE,
        NULL, write_chr_cb, NULL)
);

void bt_led_svc_set_callbacks(led_cb_t write_state,
                             led_cb_t write_blink) {
    write_state_cb = write_state;
    write_blink_cb = write_blink;
}
```

Listing 10. Plik src/bt_led_svc.c – serwis BLE

```
#pragma once

#include <zephyr/types.h>

typedef void (*led_cb_t)(unsigned int param);
void bt_led_svc_set_callbacks(led_cb_t write_state, led_cb_t write_blink);
```

Listing 11. Plik src/bt_led_svc.h – serwis BLE

```
void led_state_set(unsigned int state) {
    led_send((led_msg){ LED_STATE, state});
}
void led_blink_set(unsigned int param) {
    led_send((led_msg){ LED_BLINK, param});
}
```

Listing 12. Nowe funkcje w main.c – funkcje przekazujące parametr do sterowania LED

```
bt_control_init();
bt_led_svc_set_callbacks(led_state_set, led_blink_set);
```

Listing 13. Włączenie BLE i inicjacja serwisu LED

Bluetooth Low Energy (BLE). Dzięki temu możemy teraz ustawiać stan diody LED bezprzewodowo, korzystając np. ze smartfona.

Krzysztof Kierys
Paweł Jachimowski

Odnosniki w tekście

[1] <https://www.uuidgenerator.net/>



Sieć Thread w automatyzacji budynków

Sieci typu Mesh oparte na protokole IP zapewniają interoperacyjność, skalowalność i łatwość instalacji w aplikacjach automatyki budynkowej.

Wieżowiec Burdż Chalifa w Dubaju ma oszałamiającą wysokość 828 metrów i przewyższa drugi co do wysokości budynek na świecie, malezyjski Merdeka 118, aż o 150 metrów. Dubajski wieżowiec posiada 172000 m² powierzchni mieszkalnej i 28000 m² powierzchni biurowej, co sprawia, że nawet utrzymanie samej tylko komfortowej temperatury dla 10000 osób (bo tyle może znajdować się w budynku), stanowi spore wyzwanie. Aby mu sprostać, inżynierowie projektujący system klimatyzacji zastosowali technologię Internetu Rzeczy.

Branżowa publikacja „In-Building Tech” informuje, że zespół konserwacyjny budynku używa tysięcy czujników IoT połączonych z platformą w chmurze, aby identyfikować problemy w czasie rzeczywistym. Następnie system uruchamia inteligentne algorytmy konserwacji predykcyjnej, by uważnie monitorować system HVAC (ogrzewanie, wentylację i klimatyzację) wieżowca. Oprócz zapewnienia komfortu termicznego mieszkańcom (przy średniej temperaturze na zewnątrz wynoszącej 41°C), inteligentny monitoring systemu umożliwił 40-procentową redukcję liczby godzin potrzebnych na naprawę, przy jednoczesnej poprawie niezawodności systemu klimatyzacji do 99,95%.

Omawiany drapacz chmur jest akurat przykładem najnowocześniejszych rozwiązań w architekturze, lecz zastosowanie technologii Internetu Rzeczy może pomóc w zapewnieniu komfortu termicznego w każdym budynku – starym czy nowym. Dane

z czujników bezprzewodowych nie tylko zapewniają mieszkańcom optymalną temperaturę, ale także pomagają zmniejszyć ilość energii zużytej na ogrzewanie i chłodzenie oraz sprawiają, że systemy HVAC działają niezawodnie – a jednocześnie w przyjazny dla użytkowników sposób. Wykonanie takiego systemu może jednak stanowić wyzwanie, szczególnie na etapie projektowania i integracji w budynku, ze względu na konieczność stworzenia niezawodnych połączeń do przesyłania danych między czujnikami a chmurą.

Thread: odpowiedź na potrzeby współczesnej automatyki budynkowej

Aby zainstalować system HVAC w budynku, wymagana jest łączność pomiędzy wieloma węzłami, które będą używane do zbierania danych – tylko wtedy system centralny może automatycznie kontrolować nastawy HVAC w całym obiekcie. Przykłady takich węzłów obejmują: czujniki temperatury, wilgotności i CO₂, termostaty, czujniki obecności i systemy wentylacji uruchamiane na żądanie. Niektóre z tych węzłów są zasilane z sieci energetycznej, podczas gdy inne mają zasilane bateryjnie.

Do zagwarantowania, że budowa systemu będzie szybka, stosunkowo niedroga, a przy tym umożliwi rozbudowę i rekonfigurację, potrzebna jest bezprzewodowa technologia typu „plug-and-play”, która eliminuje okablowanie i umożliwia łatwe przenoszenie czujników, gdy zmieniają się plany urządzenia pięter. Ponadto – po zainstalowaniu sieci bezprzewodowej – ważne jest, aby można było podłączać urządzenia od dowolnego dostawcy. Dlatego też kompatybilność zastosowanych rozwiązań ma fundamentalne znaczenie dla czasu wdrożenia docelowego systemu (tzw. time-to-market).

Thread to protokół sieciowy o niskim poborze mocy i niewielkiej latencji, zbudowany przy użyciu otwartych, sprawdzonych standardów. Zapewnia on interoperacyjność między urządzeniami, dzięki obsłudze protokołu internetowego (IP) i zdefiniowanym mechanizmom „provisioningu”, które umożliwiają dodanie do sieci dowolnego zgodnego z protokołem urządzenia.

Wersja Thread 1.3.0 zawiera kilka niezwykle istotnych ulepszeń, które wprowadzają pełną funkcjonalność routingu IP oraz wykrywania usług (service discovery) do sieci Thread, a także ułatwiają jej użycie w zastosowaniach ukierunkowanych na obsługę protokołów automatyki budynkowej (na przykład KNX IoT i DALI+).

Budowa sieci Thread zazwyczaj zaczyna się od skonfigurowania routera granicznego (Thread Border Router), pod względem funkcjonalności pełni funkcję analogiczną do domowego punktu dostępowego Wi-Fi (Access Point). Główną rolą routera jest łączenie sieci Thread z innymi infrastrukturami opartymi na IP. Można go umieścić w jednym miejscu w budynku i podłączyć za pomocą Ethernet, Wi-Fi lub łączności komórkowej do sieci lokalnej i/lub do chmury. Router może pochodzić od integratora systemu HVAC lub od komercyjnego dostawcy gotowego sprzętu.

Istnieje wiele mechanizmów tzw. provisioningu (czyli mechanizmu dodawania urządzenia do sieci Thread już po zainstalowaniu routera), ale jedną z popularniejszych opcji jest użycie wstępnie udostępnionych kluczy (PSKd – Pre-Shared Key), które są dostarczane z każdym nowym produktem IoT. Skanując kod QR na urządzeniu za pomocą aplikacji na smartfonie lub dodając klucz bezpośrednio do interfejsu routera brzegowego, instalator może łatwo i w bezpieczny sposób dodać nowe urządzenie do sieci. W miarę instalacji kolejnych urządzeń, sieć automatycznie dostosowuje się i rozszerza zasięg (w zależności od roli i lokalizacji każdego węzła).

W przypadku bardziej rozległego wdrożenia w dużych budynkach, węzły zasilane z sieci energetycznej – odgrywające rolę routerów kierujących ruchem w sieci Thread – są umieszczane w kluczowych lokalizacjach, tak aby sieć mesh skalowała się, obejmując całe piętro. Taki router może być elementem systemu HVAC (na przykład termostatem) – albo prostym i niedrogim urządzeniem, którego jedynym celem pozostaje rozszerzanie sieci.

Po uruchomieniu sieci Thread można łatwo dodać inne urządzenia. Jeśli liczba wymaganych węzłów sieciowych przekroczy możliwości istniejącej instalacji, możliwe jest dodanie nowych

routerów węzłowych (zarządzających ruchem w sieci Thread) lub routerów granicznych (łączyjących sieć Thread z Internetem), aby zwiększyć zasięg sieci.

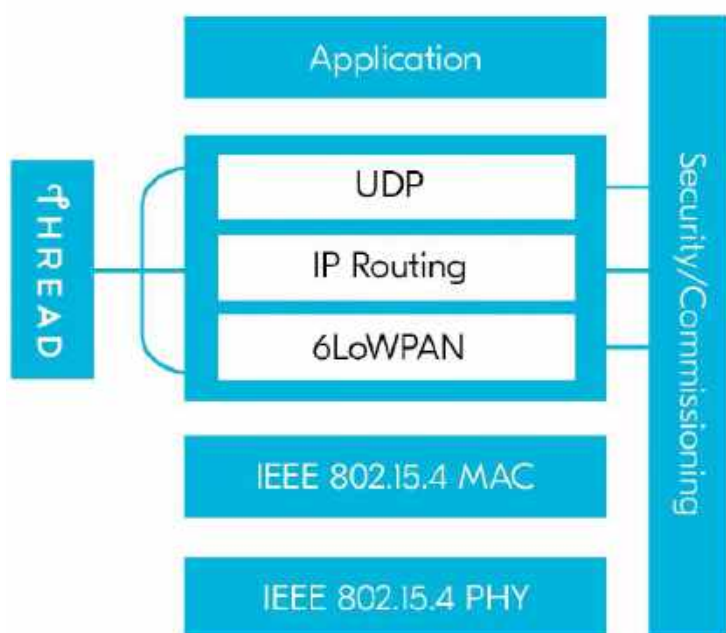
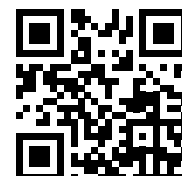
Elastyczność i interoperacyjność

Wiele węzłów w sieci Thread będzie się komunikować tylko lokalnie (na przykład czujnik temperatury wysyłający instrukcję do siłownika, aby otworzył lub zamknął zawór), podczas gdy inne będą przysyłać dane za pośrednictwem chmury. Thread obsługuje dowolną warstwę aplikacji, pod warunkiem, że używa ona protokołu IP – dzięki temu sieć jest bardzo elastyczna. Na przykład: dwa urządzenia wewnątrz budynku mogą wymieniać pomiędzy sobą dane za pomocą protokołu Matter, podczas gdy w tym samym czasie, w innym miejscu sieci inne węzły mogą się komunikować z różnymi dostawcami usług w chmurze poprzez protokół IP.

Elastyczne, interoperacyjne cechy sieci Thread odzwierciedlają cechy istniejących infrastruktur WLAN w budynkach: najpierw buduje się sieć, aby zapewnić odpowiedni zasięg, natomiast urządzenia od wielu dostawców (na przykład komputery, tablety, drukarki i kamery bezpieczeństwa) dodaje się później. Wszystkie one współistnieją bezproblemowo. Nie jest to zaskakujące – zarówno Thread, jak i Wi-Fi opierają się na technologii IP, która sprzyja interoperacyjności typu end-to-end.

Co więcej, w przypadku zastosowań przemysłowych, takich jak HVAC, Thread oferuje wbudowane, standardowe mechanizmy, umożliwiające firmom instalacyjnym szybkie projektowanie i wdrażanie sieci na dużą skalę wewnątrz budynków. Od niedawna część producentów telefonów komórkowych zaczęła implementować sprzętowe oraz programowe wsparcie protokołu Thread w swoich produktach, co dodatkowo ułatwi budowę i zarządzanie sieciami przez instalatorów. Ponieważ protokół ten jest prosty w użyciu, a przy tym oparty na IP, Thread ma szansę stać się w niedalekiej przyszłości technologią często stosowaną.

Do rozmowy o rozwiązaniach w zakresie Thread, a także Bluetooth, Wi-Fi i LTE-M/NB-IOT, firma Nordic Semiconductor zaprasza na seminarium Nordic Tech Tour, które odbędzie się już 22 października 2024 w Krakowie w hotelu Qubus. Wstęp darmowy, szczegóły i zapisy są dostępne pod adresem: <https://tiny.pl/113b1cwc>.



koktajl niusów



Przyciski MSS firmy SCHURTER Electronics z najwyższą klasą ochrony IP6K9K

Blisko 2 lata temu firma SCHURTER Electronics wprowadziła na rynek kompaktowe, a zarazem wysoce niezawodne przyciski z serii MSS. Warto podkreślić, że w chwili premiery producent niezwykle ostrożnie określił klasę ochrony IP dla przycisków z rodziny MSS jako IP67, zgodnie z normą IEC 60529. Oczywiście klasa ta okazuje się wystarczająca w zdecydowanej większości przypadków, jednakże są też i inne zastosowania, w których wymaga się wyższego niż zazwyczaj stopnia ochrony. Wiele klientów SCHURTER Electronics zapytało, czy firma może wykonać testy przycisków z serii MSS w niezależnym, certyfikowanym laboratorium, by ustalić, czy spełniają one najwyższą klasę ochrony IP6K9K. Tego typu testy są z natury dość skomplikowane, czasochłonne oraz kosztowne. Mimo to firma podjęła się tych testów, a ich wynik okazał się bardzo satysfakcjonujący. Potwierdził, że produkty te spełniają najsurowsze wymagania stawiane klasie ochrony IP6K9K. Oznacza to, że nic nie stoi na przeszkodzie, by stosować je nawet w najbardziej wymagających aplikacjach.

<https://tiny.pl/h24fn8tc>



Nowatorskie na skalę światową rozwiązanie kryptograficzne na Politechnice Warszawskiej

Naukowcy z Politechniki Warszawskiej (PW) jako pierwsi na świecie zdołali zastosować tzw. efekt motyla (teoria chaosu) do rozwiązań typu PUF. Pierwsze z wymienionych zjawisk to fenomen opisujący sytuację, w której zachodzące zmiany są zbyt złożone, żeby móc je przewidywać. Natomiast PUF (z ang. Physically Unclonable Functions) oznacza unikatowe oraz niemożliwe do skopiowania cechy fizyczne urządzeń elektronicznych, które mogą służyć jako silny mechanizm potwierdzania ich autentyczności. Punktem wyjścia

dla stworzonego na PW rozwiązania są mikroskopijne różnice w strukturze atomowej układów elektronicznych. Odpowiednie wzmocnienie tych różnic umożliwia generowanie kluczy w sposób losowy, a zarazem powtarzalny dla danego układu elektronicznego. Liczba kombinacji jest na tyle ogromna, że nie ma już potrzeby stosowania dodatkowych (skomplikowanych i kosztownych) mechanizmów ochrony haseł. Rozwiązanie opracowane przez uczonych z PW oferuje bezpieczną, a zarazem tanią i prostą w implementacji metodę generowania kluczy kryptograficznych dla różnych aplikacji. Nowy wynalazek ma szansę sprawdzić się zarówno w profesjonalnych rozwiązaniach kryptograficznych – np. do podpisu dokumentu cyfrowego – jak i w popularnych urządzeniach, w tym IoT. Będzie można z powodzeniem użyć go w najróżniejszych częściach urządzeń – od tych najmniejszych, w postaci czipów (z uwzględnieniem kart płatniczych), przez układy programowalne (np. w urządzeniach przemysłowych), po określone elementy dyskretne (zbudowane z komponentów, które razem tworzą określoną, większą całość).

<https://tiny.pl/nwcmcdp7>



Firma Orange przetestowała możliwość przeprowadzania transmisji TV za pośrednictwem łączności 5G

Na potrzeby profesjonalnej transmisji telewizyjnej firma Orange zbudowała prywatną, zamkniętą sieć 5G – w wersji „network in the box” – operującą w paśmie częstotliwości n77. Tego rodzaju sieć rezerwuje wysokie zasoby radiowe do potrzeb produkcji telewizyjnej, co okazuje się niesłychanie istotne przy obsłudze i transmisji wydarzeń z udziałem dużej liczby uczestników (takich jak koncerty czy wydarzenia sportowe). W praktyce transmisja – przeprowadzona przez Orange na częstotliwościach 3,9...4 GHz – okazała się strzałem w dziesiątkę od strony technicznej: charakteryzowała się bowiem niezwykle dużą przepustowością i małymi opóźnieniami, co jest kluczowe w przypadku produkcji telewizyjnych. Sprawdzono m.in. jakość, szybkość, a także stabilność transmisji obrazu przesyłanego na żywo z kamer Sony, wprost do wozów transmisyjnych, z użyciem ekosystemu Networked Live, który łączy hybrydowe przetwarzanie z komunikacją sieciową. Skorzystano przy tym również z kilku innych rozwiązań Sony: kodera nakamerowego HEVC, streamera CBK-RPU7 i dekodera stacjonarnego NXL-ME80.

Opisany zestaw urządzeń dostarczył niezwykle wysoką jakość obrazu przy przepustowości 15...20 Mb/s. Z kolei zaawansowany ekosystem Networked Live zapewnił elastyczność produkcji poprzez

integrację źródeł wideo i audio oraz ich synchronizację w czasie rzeczywistym, a sama transmisja udowodniła, że łączność 5G ma ogromny potencjał w zastosowaniach komercyjnych – zwłaszcza tych związanych z transmisją obrazu. Demonstracyjna wersja rozwiązania dostępna jest w tzw. 5G Labie, w warszawskiej siedzibie Orange, do której operator wraz z firmą Sony zaprasza chętnych zainteresowanych obejrzeniem opisanego rozwiązania.

<https://tiny.pl/x4cg3kkr>



Obsługa transakcji NFC oparta na mikroukładzie Secure Element – już wkrótce dostępna dla deweloperów Apple

Począwszy od wersji 18.1 systemu operacyjnego iOS, przewidziane na ten system aplikacje będą mogły obsługiwać transakcje zbliżeniowe NFC przy użyciu mikroukładu Secure Element. Już dziś za sprawą tej funkcji można m.in. dokonywać szybkich płatności w sklepach lub posługiwać się tzw. cyfrowymi kluczami do auta, natomiast w przyszłości możliwe stanie się udostępnianie za jej pośrednictwem np. dokumentów tożsamości. Na dobrą sprawę – dzięki innowacyjnemu mikroukładowi Secure Element – każdy deweloper Apple będzie mógł tworzyć bezpieczne narzędzia umożliwiające niezawodną, a zarazem godną zaufania obsługę transakcji zbliżeniowych NFC w ramach systemu iOS. Certyfikowany układ jest zgodny z branżowymi standardami i pozwala na bezpieczne przechowywanie wrażliwych informacji znajdujących się na urządzeniach mobilnych. Korzystanie z tego rodzaju oprogramowania jest maksymalnie uproszczone: po jego uruchomieniu lub skonfigurowaniu w roli domyślnej aplikacji powiązanej z funkcją NFC, do przeprowadzenia transakcji wystarczy jedynie dwukrotne naciśnięcie przycisku bocznego w iPhone. Zainteresowani programiści, którzy chcą skorzystać z możliwości mikroukładu Secure Element, muszą najpierw zawrzeć z firmą Apple umowę, a potem uiścić stosowne opłaty, by dołączyć do grona jej autoryzowanych deweloperów, spełniających określone kryteria prawne i branżowe, a także obowiązujące standardy Apple w zakresie bezpieczeństwa i prywatności – firma Apple współpracuje bowiem tylko z takimi deweloperami.

<https://nr.apple.com/dV5q9W3dE5>

Sharp Electronics świętuje 65 lat technologii fotowoltaicznej i 30 lat działalności na europejskim rynku

W 2024 roku firma Sharp Electronics świętuje aż dwie rocznice związane z działalnością w branży fotowoltaicznej: 30 lat w Europie oraz 65 lat na świecie. Jednym z zasadniczych powodów, dla których firma ta utrzymała się na rynku rozwiązań PV, jest jej pozycja jako wielkiej korporacji ze dywersyfikowaną ofertą produktów, silną marką i obecnością w wielu krajach świata. W historii Sharp Electronics można wskazać kilka zdarzeń przełomowych dla całej technologii fotowoltaicznej. W roku 1959 firma rozpoczęła prace nad ogniwami fotowoltaicznymi. Za sprawą tych działań, już



w 1970 r. ruszyła produkcja ogniw solarnych na potrzeby sektora kosmicznego. Z kolei niespełna 6 lat później przygotowano kieszonkowe ich wersje przeznaczone do kalkulatorów. Były to na tyle udane przedsięwzięcia, że przez następne lata marka Sharp Electronics notowała nieprzerwany rozwój – dość wspomnieć wyposażenie w ognia fotowoltaiczne Sharp Electronics pierwszego w historii japońskiego satelity Uma. W 2000 roku koncern cieszył się ugruntowaną już pozycją wiodącego producenta modułów fotowoltaicznych, a w 2018 roku – jako pierwszy na świecie – przygotował moduł złożony z 48 ogniw o sprawności 20%, by 2 lata później wprowadzić na rynek moduły półogniowe, a niespełna 4 lata później – również dwustronne moduły PV. Z kolei podczas targów Intersolar Europe 2024 firma zaprezentowała takie przełomowe produkty, jak np. ogniwo słoneczne z potrójnym złączem o sprawności ok. 32,65%, które z powodzeniem zastosowano jako zasilanie misji księżycowej Japońskiej Agencji Kosmicznej (JAXA). Firma Sharp Electronics wyznacza również zupełnie nowe standardy, dzięki chociażby tandemowym krzemowym ogniwom fotowoltaicznym, a także panelom, które cechuje podwójne złącze i sprawność 33,66%.

<https://tiny.pl/fhqmgzcz>



Najnowszy katalog generalny LOVATO Electric na lata 2024...2025 już dostępny

W chwili obecnej portfolio firmy LOVATO Electric liczy niemal 21 tysięcy pozycji, które podzielono na 36 rodzin produktowych. Wśród nich znajdują się m.in. nowości z dziedziny automatyki przemysłowej i zarządzania energią. Najnowszy katalog generalny LOVATO Electric na lata 2024...2025 rozpoczyna się od przeglądu nowości i ogólnej prezentacji firmy. Każdy z 36 rozdziałów poświęcony został odrębnej rodzinie produktowej. Najwięcej uwagi poświęcono przekątnikom półprzewodnikowym (SSR) i wyłącznikom kompaktowym, a także pozycjom z blokami listew rozdzielczych, przekładnikami prądowymi i urządzeniami do komunikacji zdalnej. Przewidziano również odrębny rozdział na temat bezpieczeństwa maszyn – znalazły się w nim m.in. wielofunkcyjne przełączniki, wyłączniki krańcowe i czujniki RFID. Do najważniejszych nowości produktowych należą m.in. rozłączniki izolacyjne dla fotowoltaiki, układy przełączne z napędami silnikowymi, styczniki 265...400 A oraz liczniki energii z serii DME do stacji ładowania pojazdów elektrycznych.

Cały katalog dostępny jest pod adresem: <https://tiny.pl/973y1bbj>

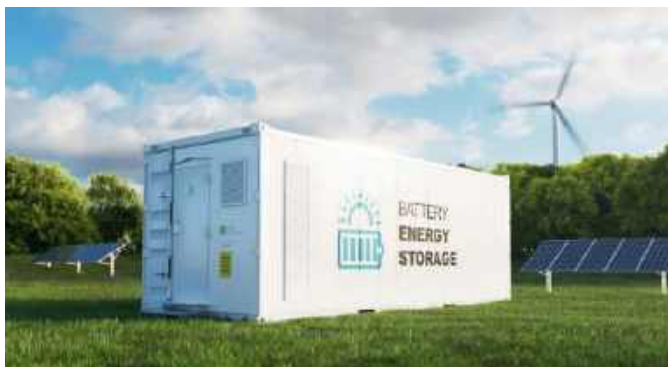
<https://tiny.pl/jkch103z>



M1000 Gen. 2 – najnowszy dysk SSD w ofercie marki GOODRAM

Ten godny uwagi następca dysku SSD stanowi zaawansowane i wydajne rozwiązanie, powstałe specjalnie z myślą o dwóch sektorach: przemysłowym i profesjonalnym. Opracowany przez markę GOODRAM dysk M1000 Gen. 2 łączy dwie istotne zalety: efektywną obsługę interfejsu SATA III oraz zgodność z formatem 2,5". W połączeniu z pamięcią NAND BiCS5 3D TLC przekładają się one na szybki i niezawodny transfer danych oraz niski pobór mocy. Szybkości odczytu i zapisu danych przez dysk wynoszą odpowiednio: 550 MB/s i 510 MB/s, co przeważnie okazuje się wystarczające w przypadku rozlicznych aplikacji przemysłowych. Produkt dostępny jest w wersjach o pojemności: 128 GB, 256 GB, 512 GB i 1 TB, pracuje w temperaturze od -40 do 85°C i stanowi świetne rozwiązanie w przypadku eksploatacji w wymagających warunkach środowiskowych. Każdy wytworzony egzemplarz dysku M1000 Gen. 2 przechodzi rygorystyczne testy jakościowe, w tym testy w komorze klimatycznej.

<https://goodram-industrial.com/aktualnosci/ssd-m1000-gen-2/>



Firma Enea zbuduje magazyn energii do stabilizacji pracy sieci

Magazyn energii o mocy 8 MW, który posłuży do stabilizacji działania sieci, a także do poprawy jakości dostarczanej energii elektrycznej, to jedna z ważnych inwestycji na terenie Wielkopolski. Firma Enea otrzymała zgodę Urzędu Regulacji Energetyki (URE) na wybudowanie baterijnego magazynu energii u zbiegu 3 ciągów linii średniego napięcia, na terenie rozdzielni sieciowej Smolnica k. miejscowości Lipka, w powiecie złotowskim w Wielkopolsce. Planowany magazyn energii – o pojemności 8 MWh – stanie się największym tego typu obiektem wybudowanym przez firmę Enea. Do jego głównych zadań należeć będą, oprócz stabilizacji mocy czynnej, przesuwanie szczytów obciążeń związanych z generowaniem i poborem energii, stabilizacja napięcia w sieciach średniego (SN) lub niskiego napięcia (NN), symetryzacja obciążenia sieci, kompensacja mocy biernej, redukcja odkształceń napięcia i wreszcie zasilanie awaryjne. Zasadniczym powodem ulokowania magazynu akurat na tym obszarze jest dynamiczny rozwój energetyki rozproszonej, co z kolei wymusza podniesienie jakości dostaw energii elektrycznej na zasilanych obszarach. Według znanych planów budowa magazynu ma potrwać do 2026 roku.

<https://tiny.pl/mskw60gt>



Wodorowe Toyoty Mirai przejechały w Berlinie – bezemisyjnie – prawie 7 milionów kilometrów

Począwszy od listopada 2022 roku flota wodorowych pojazdów Toyota Mirai wykonała niemal 550 tysięcy przejazdów ulicami Berlina, pokonując 7 milionów kilometrów – bez emisji spalin. Pojazdami tymi w szczególności przewieziono wielu gości i VIP-ów w czasie wydarzeń specjalnych – przede wszystkim Berlnale, a także GreenTechFestival, HydrogenCouncil i Federalnego Balu Prasowego. Wszystkie wodorowe Toyoty Mirai to limuzyny klasy premium, z dobrze wyposażoną kabiną mieszczącą 5 osób (wraz z kierowcą). Projektanci aut zadbali w pierwszej kolejności o nisko położony środek ciężkości, a także dość atrakcyjną stylistykę, niespotykaną nigdzie indziej. Najważniejsze zalety wodorowej Toyoty Mirai to: wynoszący do 650 km zasięg, a także 3 zbiorniki na wodór, których uzupełnienie zajmuje niespełna 5 minut. Auto napędzane jest silnikiem elektrycznym o mocy do 182 KM i momencie obrotowym 300 Nm, zasilanym ogniwami paliwowymi, które umożliwiają sprawne generację prądu na drodze reakcji wodoru z tlenem. Reakcja ta pozwala chronić środowisko za sprawą braku emisji CO₂ i innych szkodliwych substancji, a jedyne, co jest wydzielane w jej przebiegu, to para wodna. Dostępny w Toyocie Mirai napęd utrzymuje swoją sprawność i zasięg zarówno podczas mrozów sięgających -30°C, jak i przy wysokich temperaturach. Z pewnością zauważyli to pasażerowie Ubera w Berlinie: początkowo udostępniono im 50 pojazdów, aby przetestować praktyczne walory eksploatacji samochodów zasilanych wodorem w metropolii – a obecnie mają do dyspozycji flotę już 80 aut, do których wkrótce dołączy nawet 40 następnych pojazdów.

<https://tiny.pl/2jvs3wfv>

Przenośna drukarka termotransferowa Thermomark Prime 2.0 firmy Phoenix Contact o szerokim zakresie zastosowań

Kompaktowa i szybka, przenośna drukarka termotransferowa Thermomark Prime 2.0 nadaje się szczególnie do środowisk produkcyjnych – za sprawą urządzenia możliwe jest wydajne i elastyczne projektowanie przepływów pracy, co w rozległych systemach stanowi nieocenioną zaletę. Dzięki sprawdzonej technologii druku termotransferowego identyfikacja przewodów, złączy oraz kabli, a nawet urządzeń czy systemów jest precyzyjna i trwała, a znakowanie dowolnego materiału – w formie kart bądź arkuszy – zajmuje opisywanej drukarce mniej niż 8 s. Gwarantowane jest przede wszystkim proste i wygodne tworzenie znakowania do mobilnej identyfikacji za pomocą zintegrowanego oprogramowania. Wszelkie niezbędne dane wprowadzane są intuicyjnie poprzez kolorowy wyświetlacz dotykowy 7", a wymienna i wydajna bateria oraz prosta wymiana materiału i taśmy barwiącej sprawiają, że drukarka jest zawsze gotowa do pracy. Automatyczne wykrywanie taśmy barwiącej, magazynu i materiału zapobiega błędom drukowania, a z kolei do kontroli oraz zarządzania Thermomark Prime 2.0 można zastosować oprogramowanie Project Complete Marking. Gwarantuje ono duże możliwości projektowania etykiet, przejrzyste interfejsy użytkownika, a także integrację z systemami E-CAD.

<https://tiny.pl/shvrtw9r>

Jakub Tyburski
jakub.tyburski@elportal.pl

Temat numeru: Elektronika ubieralna

Urządzenia ubieralne (*wearables*) nie tak dawno były jeszcze przedmiotem zainteresowania technologicznych futurystów i autorów filmów science fiction. Miniaturyzacja układów scalonych, upowszechnienie energooszczędnych, ale wydajnych mikrokontrolerów i układów radiowych oraz postępy w dziedzinie produkcji bezpiecznych akumulatorów o zaskakująco małych wymiarach – wszystko to umożliwiło jednak rozkwit branży elektroniki noszonej na ciele. I choć ten niezwykle interesujący obszar technologii jest zdominowany przede wszystkim przez różnego rodzaju gadzety sportowe (smartwatche, inteligentne opaski czy pulsometry), to coraz częściej spotykamy się ze znacznie poważniejszymi przykładami aplikacji. Dzisiejsze urządzenia medyczne, w wielu przypadkach mające postać ubieralną, odgrywają nieocenioną rolę w monitorowaniu funkcji życiowych, zarządzaniu farmakoterapią cukrzycy czy też długookresowej diagnostyce rzadko występujących zaburzeń rytmu serca. W listopadowym Temacie Numeru przyjrzymy się najciekawszym aspektom branży aktualnego rynku urządzeń *wearable*.



Elektronika w praktyce: Wyposażenie pracowni

O aparaturze pomiarowej, sprzęcie lutowniczym czy rozmaitych narzędziach dla elektroników napisano już chyba wszystko... A jednak, na rynku ciągle pojawiają się kolejne pozycje – niektóre z nich stanowią delikatne ulepszenie już istniejących przyrządów i akcesoriów, inne całkowicie redefiniują dotychczasowy stan rzeczy i przenoszą pracę montażystów, serwisantów czy projektantów na jeszcze wyższy poziom zaawansowania. Dlatego też na łamach „Elektroniki Praktycznej” tematyka wyposażenia pracowni co pewien czas powraca. Warto być na bieżąco z najnowszymi trendami, gdyż orientacja w aktualnej ofercie rynkowej pozwala dobierać nowocześniejszy sprzęt, który nie tylko ułatwi i przyspieszy pracę, ale także zoptymalizuje ergonomię i... umożliwi poprawę uzyskiwanych efektów, chociażby w dziedzinie pomiarów czy lutowania.



Wykaz firm ogłaszających się w tym numerze „Elektroniki Praktycznej”

AKSOTRONIK.....	19
AWICAM	53
BORNICO.....	9
COMPUTER CONTROLS	7
ELMAX.....	61
HAMMOND	5
HATRON.....	55
NORDIC SEMICONDUCTOR.....	90
SEMICON.....	50, 51
TME	22, 23
WÜRTH.....	63

retroWatch

W dzisiejszych czasach – gdy postęp w dziedzinie elektroniki użytkowej nabiera rozpędu w sposób co najmniej geometryczny – coraz częściej do głosu dochodzą nuty nostalgii materializujące się w postaci urządzeń nawiązujących do stylu retro. I co ciekawe: mimo oczywistej prostoty wyżej wymienionych, jak też ich nieprzystawania do dzisiejszych możliwości technologicznych (a i zapewne części oczekiwań użytkowników), nie przeszkadza im to w zdobywaniu coraz to większej rzeszy zwolenników. Prezentowany projekt to remake kultowego swego czasu zegarka DW-2005 produkcji polskiej Unity Warel, rozslawionego przez generała Mirosława Hermaszewskiego, który wybrał ten właśnie model, aby odbyć z nim swój pierwszy (i jedyny) lot w kosmos.



Zasilacz awaryjny SuperCap UPS 5 V

Zasilacze awaryjne, znane szerzej pod akronimem UPS, pozwalają na nieprzerwaną pracę urządzeń w momencie krótkotrwałego zaniku napięcia zasilającego. Takie rozwiązania są znane od lat i stosowane ze sprzętem sieciowym. A co z tymi skromniejszymi, które zadowolają się energią dostarczaną z niewielkiego zasilacza? Do nich przewidziany jest ten właśnie projekt! Prezentowany układ podtrzymuje napięcie 5 V po zaniku zasilania z zewnętrznego źródła – rolę magazynu energii pełnią superkondensatory o długiej żywotności, zapewniające maksymalną obciążalność prądową do 300 mA i czas podtrzymania na poziomie 90 s @ 250 mA lub 6 min @ 50 mA.



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. Leszczyńska 11
tel. 22 257 84 99, e-mail: avt@avt.pl

Wydawca:
Wiesław Marciniak

Adres redakcji:
03-197 Warszawa, ul. Leszczyńska 11
e-mail: redakcja@ep.com.pl, www.ep.com.pl

Redaktor Naczelny:
Przemysław Musz

**Redaktor Programowy,
Przewodniczący Rady Programowej:**
Piotr Zbysiński

Menedżer Magazynu:
Katarzyna Gugala, tel. 22 257 84 64

Szef Pracowni Konstrukcyjnej:
Jakub Sobański

Zespół marketingu i reklamy:
Katarzyna Gugala, Bożena Krzykawska,
Grzegorz Krzykowski, Grzegorz Lalak

Stali współpracownicy:
Lucjan Bryndza, Nikodem Czechowski, Jarosław Doliński,
Andrzej Gawryluk, Krzysztof Górski, Tomasz Jabłoński,
Paweł Kowalczyk, Henryk Kowalski, Rafał Kozik,
Michał Kurzela, Jakub Nowicki, Szymon Panecki,
Adam Sobczyk, Damian Sosnowski, Ryszard Szymaniak,
Adam Tabuś, Jakub Tyburski, Robert Wołgajew

Uwaga!
Kontakt z wymienionymi osobami jest możliwy via e-mail,
według schematu: imię.nazwisko@ep.com.pl

DTP, okładka, 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

Copyright AVTKorporacja Sp. z o.o.
03-197 Warszawa, ul. Leszczyńska 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”.

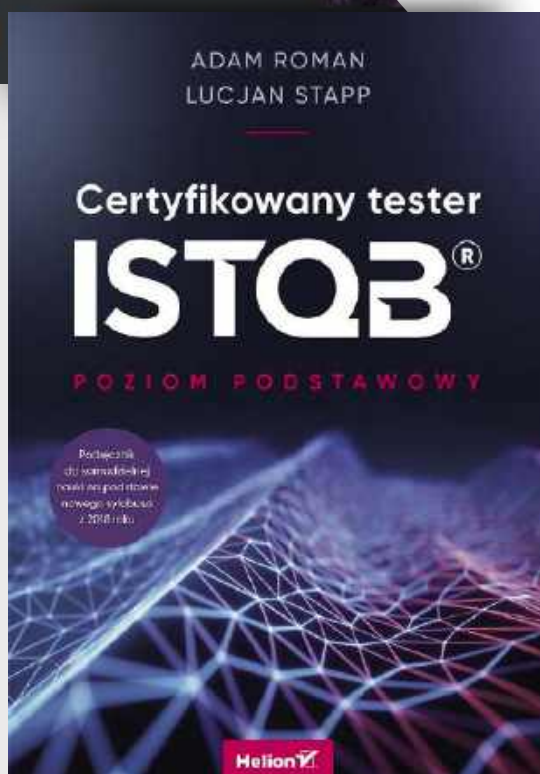


KSIĄŻKI

W ULUBIONYM KIOSKU



Z RABATEM
DO **30%**



Zobacz pełną ofertę – ponad 500 tytułów!
na www.UlubionyKiosk.pl