

ELEKTRONIKA PRAKTYCZNA

EP.com.pl

● Międzynarodowy magazyn elektroników konstruktorów ● wrzesień ● 9/2025 ●

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

- Opóźniacz wyłączenia światła • Izolowany konwerter USB-UART z translatorem poziomów
- Moduł przetwornika mocy DC w standardzie Grove

podzespoły, sprzęt, aplikacje

- Arduino R4 Nano – nie jest złe, a nawet lepiej
- Złącza do zastosowań specjalnych – wysoka niezawodność w ekstremalnych środowiskach
- Technologia Push-X: Nowa era łączenia przewodów
- Nowe silniki, przekładnie i enkodery o średnicy $\varnothing$ 16 mm – doskonała synergia zapewniająca maksymalną wydajność
- Komponenty bierne pod lupą • Wszystko, co chcieliście wiedzieć o MLCC, ale baliście się zapytać

tutoriale

- Odpadowe metody obróbki metalu • Druk 3D – moda czy niezbędne narzędzie wspierające w sytuacjach awaryjnych? • Proste sposoby na napięcie ujemne • Internet Rzeczy w pomiarach środowiskowych. Pomiar pH • Retromania, czyli jak zarobić na starociach • QuantAsylum z serii QAxix – zaawansowany system pomiarowy audio
- Syntezatory dźwięku. Informacje podstawowe

kursy

- Programowanie w środowisku MicroPython. Wyświetlacz OLED

ZŁĄCZA I AKCESORIA



KOMPONENTY BIERNE POD LUPĄ

eprasa.pl 2314d257b7



Prenumerata

Rodzina i zdrowie



-190,80 zł
114,50 zł



107,40 zł
64,50 zł



179,00 zł
107,40 zł

Dom, wnętrze



152,10 zł
91,30 zł



199,00 zł
119,40 zł



116,00 zł
69,60 zł



116,00 zł
69,60 zł

Fotografia

Zaprenumeruj
wybrane czasopisma
z **rabatem aż 40%!**

Promocja jesienna dotyczy rocznych prenumerat drukowanych.
Zamów prenumeratę na www.UlubionyKiosk.pl/prenumerata
lub poprzez dokonanie przelewu na konto:
AVT-Korporacja sp. z o.o., ul. Leszczynowa 11, 03-197 Warszawa,
ING BANK ŚLĄSKI 18 1050 1012 1000 0024 3173 1013
(w tytule wpłaty podaj nazwę czasopisma)

Muzyka i nowe technologie



220,00 zł
132,00 zł



152,90 zł
91,70 zł



-178,80 zł
107,30 zł

Masz opłaconą bieżącą prenumeratę?
Już teraz przedłuż ją z rabatem 40%.
Promocja trwa do 30.11.2025 i nie łączy się
z innymi promocjami Wydawnictwa AVT.
Koszt wysyłki na terenie kraju ponosi
wydawnictwo.

E-mail: prenumerata@avt.pl
Telefon: 22 257 84 22 (pn.-pt. 10.00–14.00)

Elektronika i automatyka



226,80 zł
136,10 zł



226,80 zł
136,10 zł



-180,00 zł
108,00 zł



179,10 zł
107,50 zł



89,40 zł
53,60 zł

Cyberbezpieczeństwo w elektronice, czyli jak zapewnić sobie spokojny sen?



Cyberbezpieczeństwo jeszcze kilka lat temu kojarzyło się głównie z komputerami, serwerami i sieciami korporacyjnymi. Dziś stanowi jeden z najważniejszych tematów w elektronice użytkowej i przemysłowej. Każdy zegarek sportowy z funkcją monitorowania zdrowia, każde urządzenie IoT w inteligentnym domu i – przede wszystkim – każde urządzenie medyczne podłączone do sieci staje się potencjalnym celem ataku. W dobie powszechnej łączności bezprzewodowej i miniaturyzacji sprzętu przestaliśmy już pytać „czy ktoś się włamie?”, a zaczęliśmy pytać „kiedy i kogo?”. To wyzwanie, które projektanci elektroniki i oprogramowania muszą traktować równie poważnie jak parametry analogowe układu czy sprawność zasilacza.

Doskonale pokazuje to głośna sprawa sprzed kilku lat, kiedy firma Medtronic musiała wycofać z rynku pompy insulinowe MiniMed z serii 508 i Paradigm. FDA ostrzegła, że w ich bezprzewodowej komunikacji występuje luka umożliwiająca osobie nieuprawnionej zdalną zmianę dawki insuliny, co mogło prowadzić do bardzo poważnych komplikacji zdrowotnych. Sama świadomość, że ktoś może przejąć kontrolę nad urządzeniem medycznym, wystarczyła, by uruchomić proces kosztownego wycofania i masowej wymiany sprzętu. Całą sprawę szeroko opisywały media, m.in. CBS News. Jeszcze bardziej dramatyczny przykład to wydarzenia z 2020 roku w Düsseldorfie, gdzie szpital uniwersytecki stał się ofiarą ataku ransomware. Systemy IT przestały działać, przez co pacjentka wymagająca pilnej pomocy została odesłana do innego miasta. Niestety jej życia nie udało się uratować. Był to prawdopodobnie pierwszy przypadek, gdy cyberatak został powiązany bezpośrednio ze śmiercią człowieka.

Te historie brzmią jak ostrzeżenia z futurystycznych thrillerów, ale wydarzyły się naprawdę i powinny być sygnałem alarmowym dla całej branży.

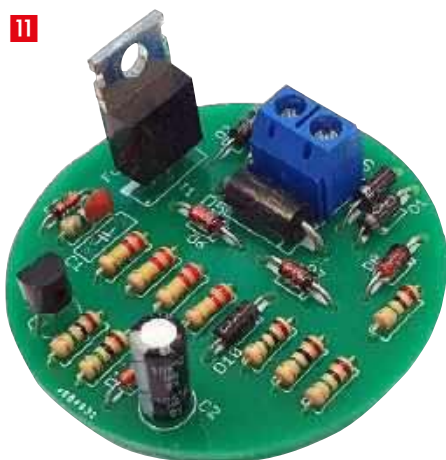
Co możemy zrobić my – elektronicy i programiści – aby nasze projekty były bezpieczniejsze? Odpowiedź nie sprowadza się już tylko do dobrych praktyk inżynierskich, choć te oczywiście zawsze są w cenie. W grudniu 2024 roku weszło w życie unijne rozporządzenie Cyber Resilience Act (CRA), a dokładniej rozporządzenie (EU) 2024/2847, które nakłada na producentów sprzętu i oprogramowania wymóg zapewnienia ciągłego bezpieczeństwa produktów cyfrowych. Od 11 września 2026 r. wejdzie w życie obowiązek raportowania podatności i incydentów, a od 11 grudnia 2027 r. pełne wymagania dotyczące certyfikacji i cyklu życia produktów. To oznacza, że temat cyberbezpieczeństwa przestaje być fakultatywny, lecz zmienia się w obowiązek prawny, z którego rozliczać nas będą nie tylko użytkownicy, ale i instytucje nadzorcze. Co więcej, od sierpnia tego roku obowiązuje norma EN 18031, która definiuje wymagania dotyczące analizy ryzyka, ochrony danych i odporności urządzeń radiowych na ataki. To kolejny dowód, że regulatorzy nie pozostawiają producentom miejsca na półśrodki.

Skoro więc wymagania rosną, a ryzyko staje się coraz bardziej namacalne, musimy myśleć o bezpieczeństwie już na etapie projektowania hardware'u. Mikrokontrolery, które jeszcze dekadę temu były prostymi jednostkami obliczeniowymi, dziś coraz częściej wyposażane są w technologie typowo „cybernetyczne”: TrustZone w procesorach ARM, pozwalający na wydzielenie stref secure i non-secure, bezpieczne obszary pamięci przeznaczone do przechowywania kluczy kryptograficznych, sprzętowe akcelerator obsługujące AES, RSA czy ECC, a także mechanizmy secure boot, które dopuszczają uruchomienie tylko prawidłowo podpisanego oprogramowania. Do tego dochodzą rozwiązania takie jak szyfrowanie pamięci Flash on-the-fly, czy dedykowane moduły root-of-trust i TPM. Jeszcze niedawno traktowane jako luksusowe dodatki, dziś stają się standardowym wyposażeniem nowych rodzin mikrokontrolerów, a projektanci urządzeń muszą nauczyć się je wykorzystywać tak samo naturalnie, jak kiedyś korzystali ze sprzętowych timerów czy przetworników A/C.

Cyberbezpieczeństwo w systemach wbudowanych to dziś ta sama „krótka koldra”, którą znamy z innych aspektów projektowania. Każde dodatkowe zabezpieczenie zwiększa złożoność projektu, zużycie zasobów i koszt urządzenia, ale jego brak może doprowadzić do czyjejś tragedii lub bankructwa producenta. Przykłady Medtronic i Düsseldorfu pokazują, że stawką jest nie tylko reputacja, ale i ludzkie życie. CRA i EN 18031 dowodzą, że regulatorzy nie zamierzają dłużej tolerować lekceważenia tych zagrożeń. A my, jako inżynierowie, musimy pogodzić się z tym, że od teraz cyberbezpieczeństwo nie jest już dodatkiem, lecz fundamentem – integralną częścią elektroniki, równie ważną jak topologia przetwornicy czy architektura układu scalonego. I tylko od nas zależy, czy tę koldrę uda się naciągnąć tak, by nie zostawić odsłoniętych żadnych krytycznych miejsc naszego systemu.

Przemysław Musz

11



Nie przeocz

Nowe podzespoły	6
Koktajl niusów	84

Projekty

Opóźniacz wyłączenia światła	11
------------------------------------	----

Miniprojekty

Izolowany konwerter USB-UART z translatorem poziomów.....	14
Moduł przetwornika mocy DC w standardzie Grove	16

Sprzęt

Arduino R4 Nano – nie jest źle, a nawet lepiej	18
--	----

Prezentacje

Druk 3D – moda czy niezbędne narzędzie wspierające w sytuacjach awaryjnych?	21
Technologia Push-X: Nowa era łączenia przewodów.....	22
Nowe silniki, przekładnie i enkodery o średnicy Ø 16 mm – doskonała synergia zapewniająca maksymalną wydajność	35

14



16



Temat numeru: Złącza i akcesoria

Złącza do zastosowań specjalnych – wysoka niezawodność w ekstremalnych środowiskach	24
--	----

Technologie wokół elektroniki

Odpadowe metody obróbki metalu	36
--------------------------------------	----

Notatnik konstruktora

Proste sposoby na napięcie ujemne	42
---	----

18



Moduły w aplikacjach

Internet Rzeczy w pomiarach środowiskowych (21). Pomiar pH.....	44
---	----

Audio bez tajemnic

QuantAsylum z serii QAxxx – zaawansowany system pomiarowy audio	50
Syntezatory dźwięku (1). Informacje podstawowe	56

44



Elektronika w praktyce

Wszystko, co chcielibście wiedzieć o MLCC, ale baliście się zapytać.....	60
Komponenty bierne pod lupą	66

Felieton

Retromania, czyli jak zarobić na starociach	72
---	----

Kursy

Programowanie w środowisku MicroPython (5). Wyświetlacz OLED	76
--	----

Prenumerata	2
Od wydawcy	3
Hity następnego numeru	87



TRZECIARĘKA ZD-11P
Uchwyt montażowy typu „Trzecia ręka”,
pająk – uchwyt z latarką, ZD11P



TRZECIARĘKA ZD-11P-1
Uchwyt montażowy typu „Trzecia ręka”,
pająk – uchwyt z latarką i lupą, ZD11P-1



TRZECIARĘKA SN-394
Uchwyt montażowy typu „Trzecia ręka”,
pająk z lupą 50 mm, przykręcany do blatu
Proskit SN-394

BESTSELLERY sklepu AVT – sklep.avt.pl

Trzecia ręka

Rabat dla Czytelników EP
przy zakupie podaj kod **EP2505TR**

-3%

Rabat dla Prenumeratorów EP
przy zakupie podaj numer prenumeraty

-6%



TRZECIARĘKA ZD-11M-1
Uchwyt montażowy typu „Trzecia ręka”,
pająk – z uchwytem na szpulkę cyny, ZD11M-1



TRZECIARĘKA ZD-11M-2
Uchwyt montażowy typu „Trzecia ręka”,
pająk – uchwyt z lupą i podświetleniem LED
ZD11M-2



TRZECIARĘKA ZD-11M-3
Uchwyt montażowy typu „Trzecia ręka”,
pająk – uchwyt z lupą i podświetleniem LED
ZD-11M-3



TRZECIARĘKA ZD-11M
Uchwyt montażowy typu „Trzecia ręka”,
pająk – uchwyt ZD11M



TRZECIARĘKA SN-392
Uchwyt montażowy typu „Trzecia ręka”
z lupą 90 mm, Proskit SN-392



TRZECIARĘKA
Uchwyt montażowy typu „Trzecia ręka”
z lupą 60 mm

nowe podzespóły

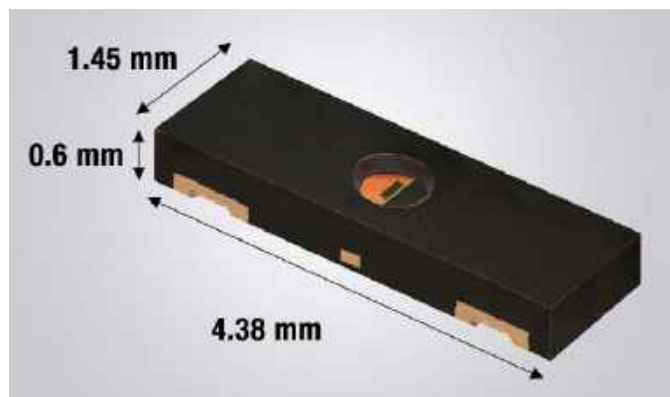
Z kilkuset nowości wybraliśmy te, których nie wolno przeoczyć. Bieżące nowości można śledzić na www.elektronikaB2B.pl



Kabel twinax Hyper Low Skew 224 Gbps o małej różnicy opóźnień sygnału

Samtec poszerza rodzinę kabli Eye Speed o kabel Hyper Low Skew Twinax, opracowany specjalnie do systemów 224 Gbps PAM4. Wyróżnia się on bardzo małymi różnicami opóźnień w parze różnicowej (skew), co jest istotne przy szybkiej transmisji sygnału. Parametr ten wynosi maksymalnie 1,75 ps/m, zapewniając bardzo dobrą stabilność przesyłu danych. Kabel może być stosowany przy częstotliwości Nyquista przekraczającej 60 GHz. Jest obecnie dostępny w wersji o przekroju 32 AWG, a w przyszłości zostanie też wprowadzona wersja 27 AWG.

www.samtec.com



Czujnik światła z kwalifikacją AEC-Q100 w obudowie o wymiarach 4,4 × 1,5 × 0,6 mm

Vishay Semiconductors powiększa ofertę komponentów optycznych o niskoprofilowy czujnik światła, zamykany w obudowie SMD o rozmiarach 4,4 × 1,5 × 0,6 mm – czyli o połowę mniejszej w porównaniu do czujników poprzedniej generacji. Model VEML4031X00 uzyskał kwalifikację AEC-Q100, dopuszczającą go do zastosowań w elektronice samochodowej. Ze względu na dużą rozdzielczość (0,0026 lx) może być umieszczany za panelami z przytłumionym szkłem. Jego czułość widmowa odpowiada charakterystyce ludzkiego oka, a dodatkowy kanał IR umożliwia rozróżnianie źródeł światła. Dzięki szerokiemu zakresowi pomiarowemu

od 0 do 172000 lx czujnik jest odporny na nasycenie podczas pracy w świetle dziennym.

Pozostałe parametry VEML4031X00:

- zakres napięcia zasilania: od 2,5 do 3,6 V,
- zakres napięcia szyny I²C: od 1,7 do 3,6 V,
- pobór prądu: typ. 0,5 μA,
- zakres temperatury roboczej: od -40 do +110°C.

www.vishay.com

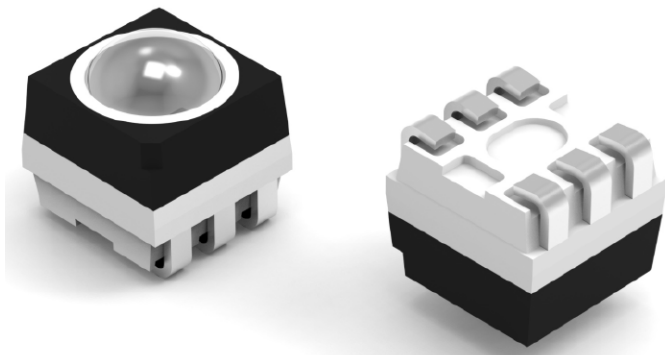


Procesory radarowe 3. generacji do systemów autonomicznej jazdy poziomej od 2+ do 4

NXP Semiconductors prezentuje nową rodzinę procesorów radarowych S32R47, przeznaczonych do pojazdów autonomicznych poziomu ADAS od 2+ do 4. Układy te zostały wykonane w technologii FinFET 16 nm. Oferują nawet dwukrotnie większą moc obliczeniową od procesorów poprzedniej generacji, a równocześnie są tańsze i wykazują mniejszy pobór mocy. W połączeniu z oferowanymi przez NXP transceiverami radarowymi mmWave, pozwalają zapewnić zgodność z systemami o poziomie nienaruszalności bezpieczeństwa ASIL D. Wewnętrzny system przetwarzania wielordzeniowego z obsługą AI/ML umożliwia lepszą klasyfikację mniejszych uczestników ruchu i obiektów (np. pieszych czy zgubionego ładunku), a także wspiera realizację bardziej zaawansowanych funkcji, takich jak DoA (Direction of Arrival).

Procesory radarowe S32R47 zawierają 4 rdzenie obliczeniowe ARM Cortex-A53 (1200 MHz) do obsługi aplikacji, 3 rdzenie ARM Cortex-M7 do wykonywania operacji czasu rzeczywistego, 8 MB pamięci SRAM oraz akceleratory radarowe (2×BBE32EP + 2×SPT3.8 do operacji FFT i 2×KQ8PPA do obliczeń zmiennoprzecinkowych). Obsługują ponadto zewnętrzne pamięci LPDDR5 i LPDDR4x. Zawierają 4 interfejsy MIPI CSI do podłączenia kamer oraz 3×2,5Gigabit Ethernet i 2×CAN FD do komunikacji z układami zewnętrznymi. Są zamykane w obudowach FCCSP o wymiarach 15×15×0,5 mm. Jako układy przeznaczone do elektroniki samochodowej, uzyskały kwalifikację AEC-Q100 i charakteryzują się szerokim zakresem temperatury roboczej struktury od -40 do +150°C.

www.nxp.com



Diody LED RGB 2,8×2,8 mm z soczewkami kopułkowymi od Würth Elektronik

Würth Elektronik poszerza ofertę diod LED RGB. W ramach serii WL-SFTD dostępne są obecnie diody produkowane w obudowach SMD typu 2828 (2,8×2,8 mm) z soczewkami kopułkowymi o kącie emisji 70° zapewniającymi równomierne oświetlenie przy minimalnych odbłaskach i maksymalnej jasności. Stopień ochrony IPx8 pozwala na zastosowania na zewnątrz pomieszczeń, a czarna, matowa obudowa zapewnia większy kontrast w porównaniu z podobnymi diodami, zamykanymi w białych obudowach PLCC.

Diody serii WL-SFTD są polecane do zastosowań w wielkoformatowych systemach projekcyjnych (video wall) oraz oświetleniu ulicznym i architektonicznym. Mogą pracować w przemysłowym zakresie temperatury. Są produkowane w technologii AlInGaP + GaN ze strukturami R, G i B o dominującej długości fali, odpowiednio, 625/470/520 nm, jasności 800/1900/400 mcd i napięciu przewodzenia 2,0/3,2/3,2 V.

www.we-online.com

Cichy przełącznik SPDT o żywotności do 2 milionów cykli

Firma Littelfuse wprowadza na rynek nowy przełącznik chwilowy TLSM o podwyższonej żywotności, mogący znaleźć zastosowanie m.in. w panelach sterujących, przyciskach wind i pilotach zdalnego sterowania. Jest to przełącznik typu SPDT o cichym przełączaniu i podwyższonej trwałości, wynoszącej 2 miliony cykli elektrycznych. Jest produkowany w obudowie SMD o powierzchni 6,1×6,0 mm i wysokości 3,45 mm. Jego parametry robocze wynoszą 1...16 V/10 μA...50 mA.

Przełącznik TLSM charakteryzuje się stopniem ochrony IP54 (świadczącym o odporności na pył i wodę) oraz szerokim zakresem temperatury otoczenia od -40 do +95°C. Rezystancja kontaktu wynosi maksymalnie 100 mΩ, rezystancja izolacji to minimum 100 MΩ, a wytrzymałość dielektryczna dochodzi do 250 VAC @ 10 mA (60 s). Układ styków SPDT daje natomiast większe możliwości konfiguracyjne w porównaniu do typowych przełączników SPST.

www.littelfuse.com



Dwukierunkowy przełącznik na bazie tranzystora GaN 650 V z podwójną bramką

Firma Infineon opracowała dwukierunkowy przełącznik IGLT65R055B2, zrealizowany na bazie 650-woltowego tranzystora CoolGaN G5 z podwójną bramką. Jest to element umożliwiający aktywne blokowanie napięcia i prądu w obu kierunkach, zrealizowany w oparciu o technologię GIT (gate injection transistor).

REKLAMA

SZKOLENIA & WEBINARY



**Certyfikowane szkolenia
Altium Designer i SOLIDWORKS**

**Webinary Altium Designer,
SOLIDWORKS i 3DEXPERIENCE**

Dowiedz się więcej na www.ccontrols.pl

COMPUTER
CONTROLS

Computer Controls Sp. z o.o.

Bielsko-Biała, ul. Bystrzańska 94

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

Może zastąpić tradycyjne przełączniki back-to-back, stosowane w konwerterach. Pracuje z maksymalnym ciągłym prądem przewodzenia 21,8 A i z maksymalnym prądem impulsowym $\pm 73,6$ A. Charakteryzuje się napięciem VSS równym 750 V, maksymalną rezystancją RSS równą 70 m Ω i ładunkiem bramki 5,4 nC.



Zintegrowanie dwóch przełączników w jednej strukturze pozwala uprościć projekty urządzeń końcowych, umożliwiając zastosowanie konwersji jednostopniowej. Zwiększa to sprawność energetyczną i niezawodność urządzenia oraz obniża koszty produkcji.

Przełącznik IGLT65R055B2 może znaleźć zastosowanie w falownikach fotowoltaicznych, systemach gromadzenia energii, stacjach ładowania i sterownikach napędów.

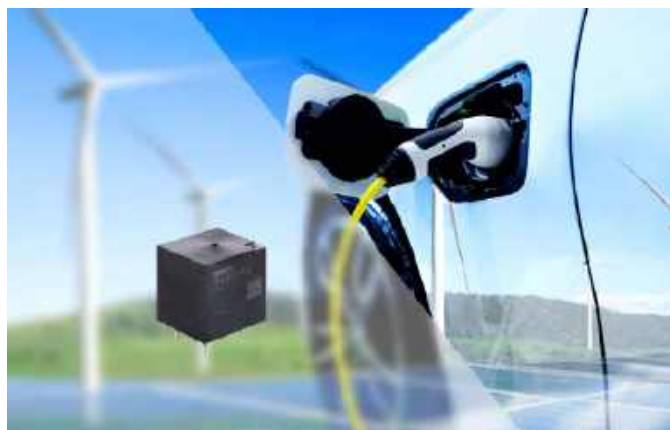
www.infineon.com

Rezystory grubowarstwowe do ochrony przed impulsami o energii do 15 J/0,1 s

Oferta komponentów zabezpieczających firmy Vishay powiększyła się w ostatnim czasie o serię rezystorów grubowarstwowych D2TO35, zamykanych w obudowach TO-263 (D²PAK). Są to rezystory o mocy znamionowej 35 W @ 25°C, umożliwiające absorpcję krótkich impulsów przepięciowych o energii do 15 J/0,1 s. Uzyskały kwalifikację AEC-Q200, pozwalającą na zastosowania w elektronice samochodowej. Poza motoryzacją mogą też znaleźć zastosowania w spawarkach, elektronarzędziach, napędach przemysłowych i aplikacjach wojskowych. Są przystosowane do pracy w temperaturze otoczenia od -55 do +175°C.



W ramach serii D2TO35H produkowane są warianty o rezystancji od 1 Ω do 14 k Ω , tolerancji od $\pm 2\%$ i rezystancji termicznej 4,28°C/W.



Małogabarytowy przekaźnik 800-woltowy do ochrony przed dużymi prądami rozruchowymi

Omron Electronics prezentuje nowy, małogabarytowy przekaźnik DC, opracowany na potrzeby ochrony przed dużymi prądami rozruchowymi w systemach ładowania o napięciu do 800 V. Model G9EJH-1-E jest zamykany w obudowie o wymiarach 31×30×27 mm. Pomimo małych gabarytów komponent spełnia wymogi normy IEC 60664 w zakresie odstępu izolacyjnego wyprowadzeń. Pracuje z napięciem cewki równym 12 V, co ułatwia jego implementowanie w elektronice samochodowej.

G9EJH-1-E to przekaźnik z kontaktami SPST NO, współpracujący z rezystorem szeregowym w układach zabezpieczających. Charakteryzuje się rezystancją kontaktów mniejszą niż 100 m Ω ,

rezystancją izolacji większą od 1 T Ω i wytrzymałością dielektryczną między cewką i kontaktami na poziomie 2500 V. Może pracować z maksymalnym prądem ciągłym 15 A i maksymalnym prądem impulsowym 30 A.

Pozostałe parametry:

- niezawodność: >200 000 cykli mechanicznych,
- czasy włączania i wyłączania: 30 ms,
- zakres temperatury roboczej: -40...+85°C.

<https://components.omron.com>

Czujnik inercyjny z dwoma akcelerometrami MEMS i wbudowaną jednostką obliczeń ML

STMicroelectronics prezentuje pierwszy na rynku czujnik inercyjny z dwoma akcelerometrami MEMS do równoczesnego śledzenia aktywności i pomiaru silnych wstrząsów. Struktura LSM6DSV320X obejmuje akcelerometry o zakresach do ± 16 g i ± 320 g oraz żyroskop o zakresie do ± 4000 dps.



Ze względu na dużą uniwersalność i małe gabaryty (3,0×2,5 mm) układ może znaleźć zastosowanie w urządzeniach mobilnych oraz aplikacjach smart home, smart industry i smart driving. Przykładem mogą być trackery rejestrujące trening sportowców, kontrolery do gier, sprzęt ochrony osobistej używany w przemyśle czy też monitorowanie konstrukcji. Poza wymienionymi wcześniej blokami układ oferuje jednostkę MLC (machine-learning core), czyli procesor AI przeprowadzający wnioskowanie bez potrzeby komunikowania się z układami zewnętrznymi. Pozwala on obniżyć pobór mocy i przyspieszyć działanie aplikacji. Wszystkie wewnętrzne czujniki są wzajemnie zsynchronizowane. Podobnie jak w przypadku pozostałych czujników smart MEMS firmy STMicroelectronics, LSM6DSV320X oferuje funkcję adaptacyjnej autokonfiguracji (ASC), pozwalającą zredukować pobór mocy. Czujniki z ASC mogą automatycznie dostrajać parametry pracy w czasie rzeczywistym po wykryciu określonego wzorca ruchu lub sygnału z MLC, bez konieczności interwencji mikroprocesora nadrzędnego.

Pozostałe właściwości:

- interfejsy SPI/I²C i MIPI I³C v1.1,
- do 4,5 kB pamięci Smart FIFO,
- napięcie zasilania od 1,62 do 3,6 V,
- pobór prądu do 0,8 mA,
- obudowa LGA-14L (3,0×2,5×0,83 mm),
- zakres temperatury roboczej od -40 do +85°C.

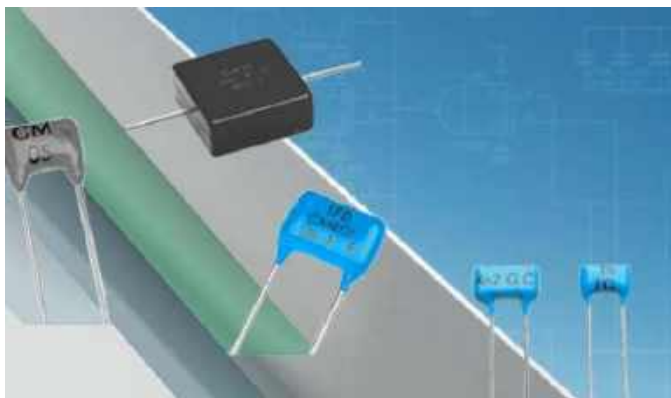
www.st.com

Potencjometry do gitar elektrycznych i profesjonalnego sprzętu audio

W ofercie firmy Bourns pojawiła się nowa seria potencjometrów PDB241-GTR, zaprojektowanych specjalnie do zastosowań w gitarach elektrycznych i profesjonalnym sprzęcie audio. Wyróżniają się one długą żywotnością na poziomie 100 000 cykli mechanicznych i małym momentem obrotowym. Występują w wariantach o rezystancji od 10 k Ω do 1 M Ω , charakterystyce liniowej lub logarytmicznej oraz z gładkim lub radełkowanym zakończeniem wałka. Mogą pracować w temperaturze otoczenia od -10 do +70°C. Oferta obejmuje obecnie dwa modele, różniące się długością wałka, wynoszącą 19 mm w przypadku PDB241-GTR01 oraz 27,5 mm w przypadku PDB241-GTR03.



www.bourns.com



Kondensatory mikowe do aplikacji wojskowych i lotniczych

Firma Exxelia dodała do oferty cztery serie kondensatorów mikowych wysokiej klasy do zastosowań wojskowych, lotniczych, medycznych i pomiarowych. Wyróżniają się one bardzo dużą stabilnością i niezawodnością w funkcji czasu i temperatury. Ponadto wykazują małe straty, co jest istotne w aplikacjach w.c.z. Mogą pracować z napięciem roboczym do 1000 VDC.

Seria CMR obejmuje precyzyjne kondensatory wysokiej niezawodności, zgodne z wymogami standardu MIL-PRF-39001 i mogące znaleźć zastosowania m.in. w obwodach RF. Charakteryzują się one dużą dobrocią i są dostępne w wersjach o pojemności z zakresu od 1 pF do 12 nF oraz napięciu znamionowym do 500 V. Mogą pracować w temperaturze otoczenia od -55 do +150°C.

W ramach serii CA dostępne są hermetyczne kondensatory do układów, w których kluczowa jest niezawodność. Wykazują dużą dobroć i małą rezystancję ESR. Ich zakres pojemności rozciąga się od 5 pF do 100 nF, a napięcie znamionowe może wynosić od 63 do 1000 V. Zakres temperatury roboczej to -55...+125°C.

Seria	Dielektryk	Pojemność	Napięcie znamionowe	Tolerancja
CMR	mika	1 pF...12 000 pF	do 500 V	±2%
CA		5 pF...100 nF	do 1 kV	±2%
MF		4,7 pF...33 nF	do 1 kV	±2%
CM		10 pF...1 µF	do 1,5 kV	±1%, ±2%, ±5%, ±10%

Precyzyjne kondensatory dużej niezawodności serii MF są polecane do układów w.c.z. i czułych przyrządów pomiarowych oraz ogólnie do aplikacji o znaczeniu krytycznym. Do ich zalet należy duża dobroć i mały współczynnik temperaturowy. Dostępne są warianty o pojemności od 4,7 pF do 33 nF i napięciu znamionowym do 1 kV.

Seria CM obejmuje kondensatory o pojemności od 200 pF do 12 nF i napięciu znamionowym od 100 do 500 V, mogące pracować w zakresie temperatury otoczenia do +150°C. W zakresie niezawodności i odporności na narażenia środowiskowe spełniają one wymogi normy MIL-C-5. Małe straty czynią je idealnymi do zastosowań w aplikacjach wojskowych i lotniczych, w których kluczowa jest odporność na ciężkie warunki środowiskowe. Kondensatory CM są dostępne w wersjach o tolerancji już od ±1%.

www.exxelia.com

Czujniki światła o małych szumach i dużej czułości

Firma Hamamatsu Photonics wprowadza na rynek serię krzemowych czujników światła S17353 o małych szumach i dużej czułości, przeznaczonych do pracy na długości fali do 800 nm. Są to elementy nadające się idealnie do przyrządów analitycznych, pracujących w zakresie małego natężenia światła.

W ramach serii S17353 dostępnych jest 6 typów czujników, mogących znaleźć zastosowanie m.in. w spektroskopii, analizie



	Obudowa	Powierzchnia światłoczuła [mm]	Zakres widmowy [nm]	Czułość [A/W]	Prąd ciemny [nA]
S17353-02K	TO-5	Ø 0,2	320...1000	typ. 0,44	1
S17353-05K		Ø 0,5			1,5
S17353-10K		Ø 1,0			2
S17353-20K		Ø 2,0			5
S17353-30K	TO-8	Ø 3,0			10
S17353-50K		Ø 5,0			25

REKLAMA



- Przewody
- Oploty
- Wiązki

Semicon Sp. z o.o.

ul. Zwoleńska 43/43a, 04-761 Warszawa, 22 615-73-71

semicon.com.pl

jszyszko@semicon.com.pl



- Sygnałowe
- Światłowodowe
- Koncentryczne
- Wysokonapięciowe



- Przemysł i automatyka
- Medycyna
- Lotnictwo i obronność
- Broadcast i audio-wideo



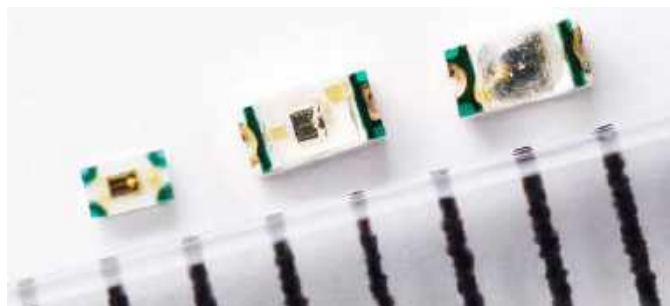
- Wielożyłowe
- Hybrydowe
- Nawojowe
- Miedziane oploty

Innowacyjne produkty
Innowacyjne technologie



chemicznej i przemyśle. Są one produkowane w obudowach TO-5 i TO-8. Charakteryzują się czułością wynoszącą typowo 0,44 A/W, prądem ciemnym od 1 do 25 nA i średnicą powierzchni światłoczułej od $\varnothing$ 0,2 mm do $\varnothing$ 5,0 mm.

www.hamamatsu.com



Diody LED NIR o dużym natężeniu promieniowania do aplikacji AR/VR

Firma Rohm dodaje do oferty diod LED NIR nowe wersje o dużym natężeniu promieniowania, zaprojektowane do zastosowań w przemysłowych czujnikach optycznych, drukarkach, pulsoksymetrach oraz aplikacjach AR/VR (śledzenie wzroku, rozpoznawanie gestów). Są one dostępne w 6 wariantach, produkowanych w 3 typach obudów. Dwa nadajniki, stanowiące rozszerzenie rodziny Picoled – SML-P14RW i SML-P14R3W – są zamykane w miniaturowych obudowach o wymiarach 1,0×0,8×0,2 mm. Cztery pozostałe są dostępne w standardowych obudowach o powierzchni 1,6×0,8 mm. CSL0902RT i CSL0902R3T zawierają okrągłe soczewki o wąskim kącie emisji, natomiast CSL1002RT i CSL1002R3T to diody z soczewkami płaskimi o szerokim kącie emisji. Każdy z wariantów obudowy jest dostępny w dwóch wersjach, pracujących na długości fali 850 nm lub 940 nm.

	Wartości maksymalne			Charakterystyka elektro-optyczna			Obudowa [mm]
	IF [mA]	VR [V]	Zakres temp. pracy [°C]	IE [mW/sr]	VF [V]	λp [nm]	
SML-P14RW	50	5	-40...+85	2,8	1,5	860	1,0×0,6×0,2
SML-P14R3W						940	
CSL0902RT	30			5,0	1,4	850	1,6×0,8×1,24
CSL0902R3T						940	
CSL1002RT	30			1,9	1,4	850	1,6×0,8×1,06
CSL1002R3T						940	

Diody 850 nm są idealne do aplikacji AR/VR o dużej czułości, np. do śledzenia wzroku i wykrywania gestów. Diody 940 nm, mniej podatne na działanie światła słonecznego, nadają się do zastosowań w czujnikach ruchu. Mogą być również stosowane m.in. w czujnikach pulsoksymetrycznych.

www.rohm.com

Seria miniaturowych odbiorników GNSS o powierzchni 23×16 mm

Firma Septentrio wprowadza na rynek serię miniaturowych odbiorników GNSS o oznaczeniach mosaic-G5, dostępnych w trzech wariantach: P1, P3 i P3H. Uzupełniają one wcześniejszą ofertę, obejmującą odbiorniki mosaic-X5, mosaic-H i mosaic-T, w stosunku do których oferują o 60% mniejszą powierzchnię, a także o 40% niższy pobór mocy. Są zamykane w obudowach o powierzchni 23×16 mm i masie jedynie 2,2 g, dzięki czemu doskonale nadają się do zastosowań w dronach i robotyce.

mosaic-G5 P1 to odbiornik 3-zakresowy, mosaic-G5 P3 obsługuje 4 zakresy, natomiast mosaic-G5 P3H – 3 zakresy. Ostatni z wymienionych modułów to rozwiązanie typu „heading”, obsługujące pracę z dwiema antenami GNSS jednocześnie i potrafiące wyznaczyć wektor kierunku na podstawie różnicy faz odbieranych sygnałów GNSS.



Wszystkie trzy modele obsługują konstelacje GPS/GLONASS/Beidou/Galileo/QZSS i pozycjonowanie różnicowe RTK, oferują ponadto opracowany przez Septentrio tryb antyzakłóceń AIM+.

Pozostałe cechy:

- dokładność pozycjonowania w poziomie/pionie (DGNSS): 0,4 m/0,7 m,
- dokładność pozycjonowania w poziomie/pionie (SBAS): 0,6 m/0,8 m,
- dokładność pozycjonowania w poziomie/pionie (tryb autonomiczny): 1,2 m/1,9 m,
- interfejs USB 2.0 device,
- czas inicjalizacji równy 7 s.

Obecnie odbiorniki mosaic-G5 są dostępne w wersjach próbnych, a uruchomienie masowej produkcji ma nastąpić jeszcze w bieżącym roku.

www.septentrio.com



Wysokoprądowe złącza zasilające DC do aplikacji EV

Do oferty firmy TME wchodzi złącza zasilające PowerLok produkcji Amphenol GEC, zaprojektowane do sektorów: motoryzacyjnego oraz energetycznego (źródła OZE, magazyny energii). Oferta obejmuje wtyki i gniazda PowerLok G2 oraz złącza PowerLok 4.0 G2.

Wtyki i gniazda PowerLok G2 zostały wykonane ze stopu aluminium i charakteryzują się stopniem ochrony IP67. Mogą przewodzić prądy o natężeniu do 300 A przy napięciu znamionowym sięgającym 1 kVDC. Występują w wielu konfiguracjach, w tym prostej i kątowej, zawierających od 1 do 3 pinów. Oferta obejmuje wersje do mocowania panelowego oraz na przewodach. Konstrukcja korpusu zapewnia pełną izolację i ekranowanie styków w celu ochrony operatorów i minimalizacji generowanych zaburzeń elektromagnetycznych.

Złącza PowerLok 4.0 G2 są przeznaczone do zasilania akcesoriów i podzespołów w pojazdach, w tym czujników, siłowników i serwomechanizmów. Mogą one przewodzić prądy o natężeniu do 60 A. Po połączeniu charakteryzują się stopniem ochrony IP69. Ich zakres temperatury roboczej rozciąga się od -40 do +125°C.

www.tme.eu



W ofercie AVT*

AVT6086

Najważniejsze parametry:

- opóźnienie wyłączenia obciążenia o około 50...90 sekund,
- prosta instalacja: układ podłącza się równolegle do styków przelaznika sieciowego,
- możliwość instalacji modulu w typowej puszcze podtynkowej,
- układ przystosowany do napiecia przemienneo 230 V/50 Hz,
- współpraca z obciażeniem o mocy z przedziału 2...30 W, najlepiej lampą LED.

* **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] – płytka drukowana bez elementów i dokumentacji.
- Kity, w których występuje układ scalony wymagający zaprogramowania, mają następujące dodatkowe wersje:
- wersja [A*] – płytka drukowana [A] + zaprogramowany układ [UK] i dokumentacja,
 - wersja [UK] – zaprogramowany układ.

Projekty pokrewne na stronie www.ep.com.pl

- (aktywne linki do artykułów):
- Zdalny włącznik radiowy
 - Wyłącznik zasilania z opóźnieniem
 - Bezprzewodowy wyłącznik wi-link
 - Wyłącznik czasowy ustawiany enkoderem

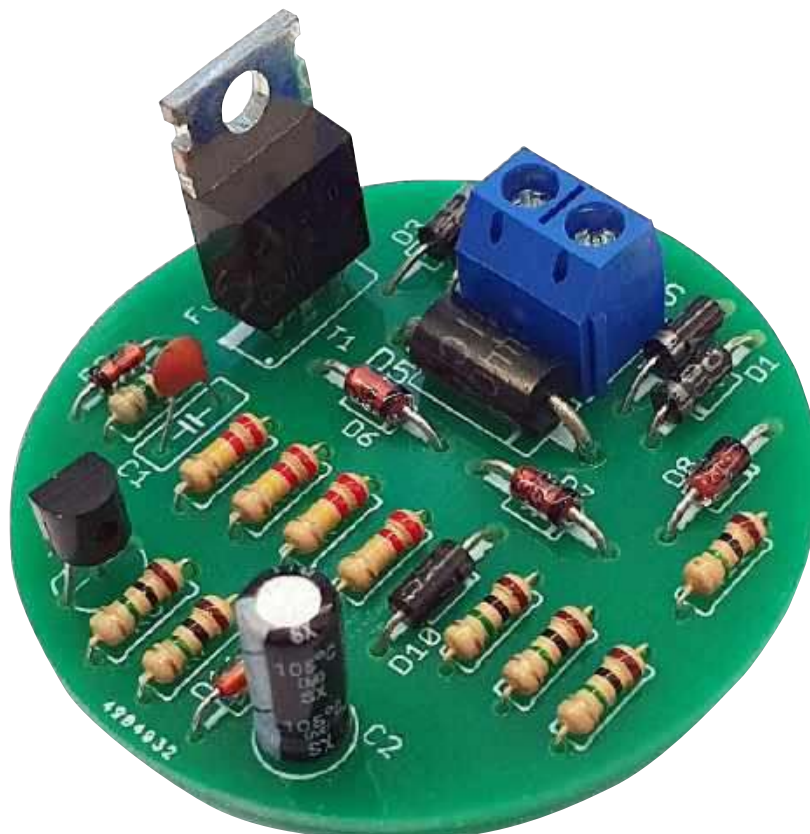
Nie każdy zestaw AVT występuje we wszystkich wersjach! Każda wersja ma załączony ten sam plik PDF! Podczas składania zamówienia upewnij się, którą wersję zamawiasz! <http://sklep.avt.pl>

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

Opóźniacz wyłączenia światła

Jak działa wyłącznik światła, każdy doskonale wie: w jednej pozycji klawisza światło świeci, w drugiej – nie. Jeden ruch ręką pozwala natychmiastowo zmienić bieżący stan na przeciwny. Niekiedy jednak przydałoby się wprowadzić pewne opóźnienie wyłączenia źródła światła, aby można było np. włożyć klucz do zamka i go swobodnie przekręcić, bez robienia tego po omacku lub przy świetle latarki trzymanej w zębach. Ten prosty i tani układ umożliwi wprowadzenie takiej funkcji zwłocznej do nowoczesnych źródeł światła.

Prezentowany układ należy zaliczyć do szerokiego grona prostych urządzeń, które ułatwiają nam codzienne życie. Tu o czymś przypomnę, tam coś za nas załączę, gdzieś indziej zrobią coś automatycznie. Opisane urządzenie – mimo że jest proste w działaniu – również upraszcza nam funkcjonowanie. Można wyjść z ciemnego korytarza, „pacnąć” naścienne wyłącznik dłońią i bez obaw przejść do następnego pomieszczenia, które jeszcze nie jest oświetlone, by podobnym



gestem włączyć źródło jakże potrzebnego naszym oczom światła. Blask dochodzący z poprzedniego, opuszczonego już pomieszczenia, będzie nas w tym wspomagał, po czym samoczynnie zgaśnie – nie ma więc potrzeby wracania do korytarza, który już opuściliśmy.

Projektując ten układ, miałem na uwadze przystosowanie go do współczesnych realiów, w których lwia część źródeł światła stanowią lampy LED o mocy nieprzekraczającej kilkunastu watów. Nie trzeba zatem wymieniać oświetlenia na żarowe, aby tylko triak (lub inny element wykonawczy) mógł pracować w prawidłowych warunkach. Dodatkowym atutem jest bardzo prosta instalacja.

Budowa układu

Niekiedy siła tkwi w prostocie. Prezentowany w artykule układ nie zawiera

żadnego mikrokontrolera ani innego układu scalonego. Schemat ideowy urządzenia znajduje się na **rysunku 1**. Zaciski złącza J1

REKLAMA

LASEROWE SZABLONY DO MONTAŻU SMT

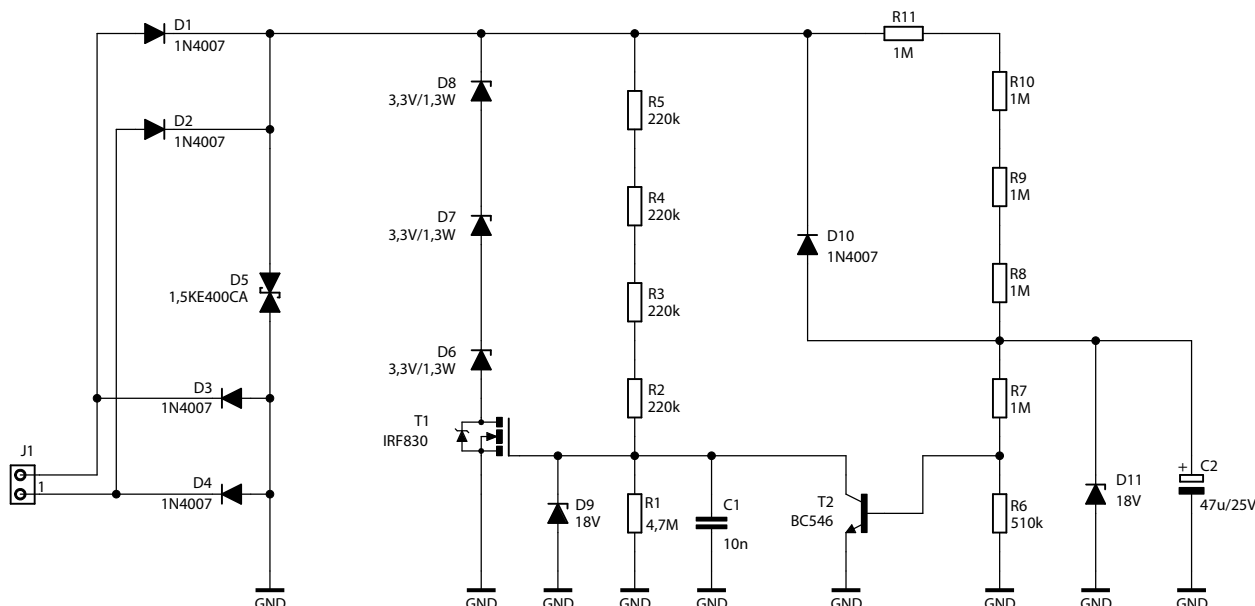
Materiał: stal nierdzewna CrNi
Zakres grubości blach: 0,020–1,000 mm
Wycinamy również detale o dowolnych kształtach



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

WYSOKIE NAPIĘCIE

W układzie występuje wysokie napięcie, które może spowodować porażenie elektryczne. Układ może być zmontowany i podłączony tylko przez osoby wykwalifikowane, posiadające odpowiednie uprawnienia. Zdecydowanie nie zalecamy budowy układu osobom pocztującym!



Rysunek 1. Schemat ideowy opóźniacza wyłączenia światła

przewodzą do styków przełącznika, więc napięcie na nich będzie równe sieciowemu (przy rozwarciu tychże) lub zerowe po ich zwarcu przez użytkownika.

Elementem wykonawczym, który nie ma szczególnych wymagań co do minimalnego prądu przewodzenia, jest tranzystor MOSFET. Jednak przewodzi on jednokierunkowo, tymczasem prąd w naszych sieciach elektroenergetycznych ma charakter przemienny. Do ominięcia tego problemu wykorzystano starą sztuczkę, która królowała w czasach, kiedy tyrystory były na wagę złota, zaś o triakach nikt jeszcze nie słyszał. Otóż prąd zasilający odbiornik (tutaj: lampę LED) przepływa przez mostek Graetza, który czyni go jednokierunkowym po stronie elementu wykonawczego. Proste, tanie i skuteczne, zważywszy na to, że prąd płynący przez diody mostka Graetza jest niewielki, toteż moc w nich wydzielana jest pomijalnie niska. Dioda D5 zabezpiecza tranzystor wykonawczy przed uszkodzeniem w razie wystąpienia impulsu wysokiego napięcia, na przykład podczas przełączania obciążenia o charakterze indukcyjnym.

Wysokonapięciowy tranzystor T1, wraz z czterema diodami D1...D4, służą do „zastąpienia” styków wyłącznika wtedy, kiedy zostanie on rozarty. Jest jednak pewien drobny mankament: po załączeniu tranzystora MOSFET, spadek napięcia na nim wyniesie kilkadziesiąt miliwoltów lub niewiele więcej, skąd zatem wziąć zasilanie niezbędne dla podtrzymania potencjału jego bramki? W tym celu zostały dodane trzy

połączone szeregowo diody Zenera, na których odkłada się napięcie o łącznej wartości około 10 V. Tyle wystarczy do spolaryzowania bramki tranzystora unipolarnego i utrzymania go w stanie przewodzenia. Niestety, o tyle mniej będzie wynosiło napięcie trafiające na odbiornik, lecz współczesne lampy LED mają na ogół wbudowaną prostą przetwornicę, która doskonale radzi sobie z niezmiennie niższym napięciem zasilającym.

Bramka tranzystora MOSFET wymaga napięcia stałego o wartości nieprzekraczającej ± 20 V względem jego źródła, zatem tuż przy tranzystorze znalazła się dioda Zenera o napięciu przebicia 18 V. Jej prąd ograniczają połączone szeregowo rezystory R2...R5. Użycie kilku elementów szeregowych pozwala rozłożyć napięcie, więc nie występuje ryzyko przebicia rezystorów – dotyczy to sytuacji przy rozwarciu styków wyłącznika, bowiem odłoży się na nich niemal całe napięcie sieciowe. Niemal całe, bowiem jego część pozostanie na wspomnianej już diodzie Zenera D9. Kondensator C1 o relatywnie niewielkiej pojemności stanowi prosty filtr tętnień napięcia, które występują na wyjściu prostownika dwupołkowego. Rezystor R1 rozładowuje C1 oraz pojemność bramki-źródła tranzystora T1.

Opisany wyżej obwód potrafi podtrzymać świecenie lampy niezwłocznie po rozwarciu styków wyłącznika ściennego. Potrzebna jest jeszcze jedna część, która wyłączy tranzystor T1 po pewnym czasie od owego rozwarcia. Do tego służy tranzystor T2, który wchodzi w nasycenie i zwi

erza bramkę ze źródłem tranzystora T1, skutecznie go zatykając. Jednak musi on się załączyć z opóźnieniem, które zapewnia kondensator C2 – ładujący się powoli przez szeregowo połączone rezystory R8...R11. Podobnie jak w przypadku wcześniej omówionych R2...R5, tak i te elementy zostały rozdzielone w celu rozłożenia wysokiego napięcia.

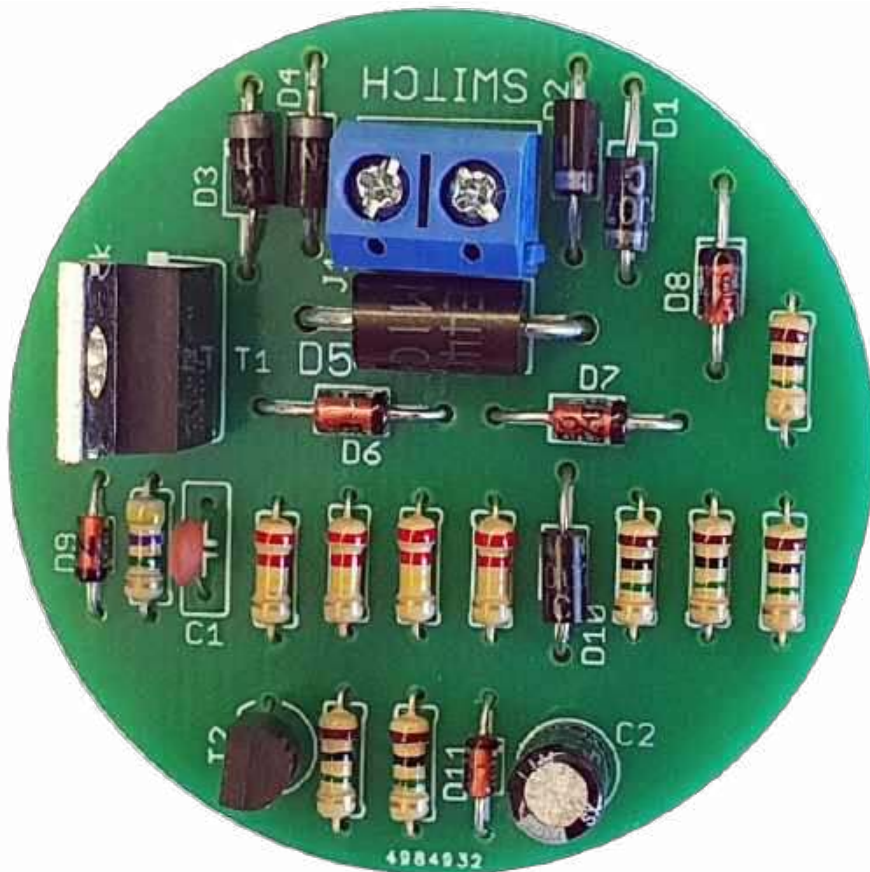
Kondensator C2 ładuje się w miarę upływu czasu (od rozłączenia styków wyłącznika), wraz z tym rośnie potencjał bazy tranzystora T2. Napięcie z C2 jest dodatkowo dzielone przy użyciu rezystorów R6 i R7, aby załączenie T2 nie nastąpiło zbyt szybko. Dodatkowo, rezystory te rozładowują C2 po zwarcu styków wyłącznika sieciowego. Dioda D11 ogranicza maksymalne napięcie, do jakiego może naładować się C2, aby nie uległ on uszkodzeniu. Dioda D10 przyspiesza rozładowywanie C2 po zwarcu styków włącznika.

Montaż i uruchomienie

Układ został zmontowany na niewielkiej, jednostronnej płytce drukowanej o okrągłym obrysie. Jej średnica wynosi 48 mm, więc łatwo można ją zamontować pod już istniejącym wyłącznikiem podtynkowym, ponieważ PCB zmieści się w każdej puszcze elektrycznej. Wzór ścieżek oraz schemat montażowy płytki pokazano na **rysunku 2**. Nie przewidziano na jej powierzchni żadnych otworów montażowych – jest na tyle mała i lekka, że utrzyma się zawieszona na przewodach znajdujących się za wyłącznikiem.

REKLAMA

facebook.com/ElektronikaPraktyczna



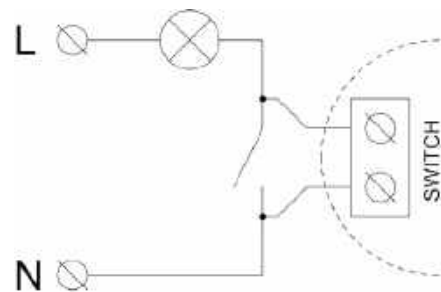
Fotografia 1. Widok zmontowanego układu

Montaż proponuję przeprowadzić w typowy sposób, czyli zaczynając od elementów leżących płasko na powierzchni laminatu, jak diody i rezystory. Na sam koniec proponuję zostawić tranzystor T1, ponieważ jest on najwyższym podzespołem, więc może utrudniać montaż, gdyby znalazł się na płytce we wcześniejszym etapie. Gotowy do działania układ można zobaczyć na fotografii 1. W roli T1 może pracować dowolny inny tranzystor wysokonapięciowy MOSFET-N w obudowie TO220, o napięciu dren-źródło nie mniejszym niż 500 V oraz o rezystancji otwartego kanału nieprzekraczającej kilku omów – pozostałe parametry nie są krytyczne w tym układzie.

Ze względu na wymiary oraz sposób podłączenia modułu, istniejąca instalacja elektryczna nie wymaga jakichkolwiek

przeróbek. **Rysunek 3** zawiera schemat podłączenia. W momencie działania układu, czyli przez kilkadziesiąt sekund po rozłączeniu styków przełącznika, na obciążenie trafia napięcie niższe o ok. 10 V od napięcia sieciowego. Po upływie tego czasu sterowana lampa LED gaśnie płynnie w przeciągu około 2 s.

Należy pamiętać, iż przez opisywany układ przepływa pewien niewielki prąd w stanie wyłączenia lampy, czyli przy rozwarciu styków wyłącznika – jego wartość szczytowa to kilkaset nanoamperów. Aby kondensator C2 prawidłowo się rozładował przed następnym cyklem pracy, wyłącznik



Rysunek 3. Schemat podłączenia układu do istniejącej instalacji

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

R1: 4,7 MΩ

R2...R5: 220 kΩ

R6: 510 kΩ

R7...R11: 1 MΩ

Kondensatory:

C1: 10 nF MKT (raster 5 mm)

C2: 47 μF/25 V (raster 2,5 mm)

Półprzewodniki:

D1...D4, D10: 1N4007

D5: 1,5KE400CA

D6...D8: dioda Zenera 3,3 V (1,3 W)

D9, D11: dioda Zenera 18 V (0,5 W)

T1: IRF830 (opis w tekście)

T2: BC546

Pozostałe:

J1: złącze śrubowe ARK2/500

musi pozostać zwarty przez minimum 1 minutę. Jeżeli zostanie rozarty wcześniej, ustalony czas może okazać się krótszy. Wartość opóźnienia rzędu 90 s należy traktować jako orientacyjną, bowiem silny wpływ na ten parametr mają zarówno rozruty elementów, moc pobierana przez lampę, jak i temperatura otoczenia. Podczas testów okazało się, że podłączenie lampy o mocy 2 W daje czas opóźnienia rzędu 90 s, zaś przy lampie o mocy 20 W wartość ta zmniejszała się do około 50 s z uwagi na wyższe napięcie odkładające się na diodach Zenera – bowiem płynął przez nie prąd o wyższym natężeniu. To z kolei przyspieszało ładowanie C2.

Michał Kurzela, EP

REKLAMA

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

WWW.AKSOTRONIK.COM.PL

Magnesy neodymowe oraz ferrytowe
Ceny od 6,30zł

Przełączniki klawiszowe wodoodporne/płaskie
Ceny od 2,40zł

Przewodniki do przewodów
Ceny od 11,00zł

Druty oporne od 0,15 do 0,8 mm
Ceny od 5,70zł

Kaski elektryczne żarłokowe
Ceny od 0,32zł

Szeroki węgiel do elektronizacji
Ceny od 2,60zł

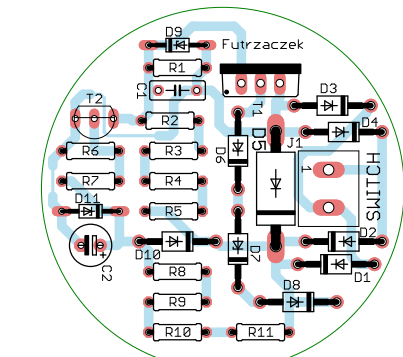
Przełączniki do elektronizacji zwykłe i elektromagnetyczne
Ceny od 7,00zł

Podłoża/organizery
Ceny od 0,95zł

Złącza termiczne Superseal
Ceny od 1,10zł

Zestawy diod M2, M3 z radiatorami i podkładkami
Ceny od 2,50zł

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



Rysunek 2. Schemat montażowy i wzór ścieżek płytki



W ofercie AVT*

AVT6087

Najważniejsze parametry:

- obsługa poziomów logicznych 1,2...5 V,
- zintegrowana przetwornica DC/DC zasilająca konwerter poziomów,
- pobór prądu z urządzenia zewnętrznego (po stronie portu UART): 20 µA (maks.)/2...6 µA (typ.),
- kompatybilność ze standardem USB 2.0,
- wbudowane diody LED sygnalizujące obecność napięć zasilania i kierunku transmisji,
- popularny układ FT232RL, obsługiwany przez większość systemów operacyjnych.

* **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 wzlutować w dotychczasową płytkę drukowaną (PCB). Wykaz elementów znajduje się w dokumentacji, która jest podlinkowana w opisie kitu. Mając na uwadze różne potrzeby naszych klientów, oferujemy dodatkowe wersje:

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

Projekty pokrewne na stronie www.ep.com.pl

(aktywne linki do artykułów):

- Konwerter USB-UART z ekstenderem
- Dwukanałowy konwerter USB-C z układem FT232H
- Konwerter USB-C-RS485
- Konwerter USB-UART w standardzie Grove
- Konwerter USB/RS232 z izolacją galwaniczną
- Konwerter KF dla tunera RTL-SDR
- Miniaturowy konwerter USB-UART
- Miniaturowy konwerter USB-RS485

Nie każdy zestaw AVT występuje we wszystkich wersjach! Każda wersja ma załączony ten sam plik PDF! Podczas składania zamówienia upewnij się, którą wersję zamawiasz!

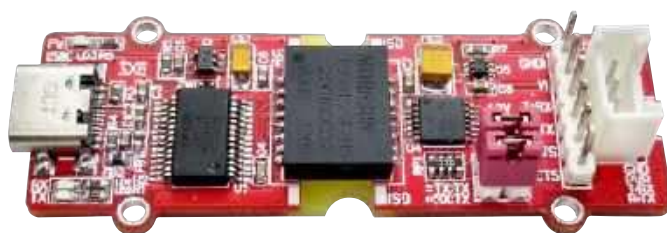
<http://sklep.avt.pl>

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

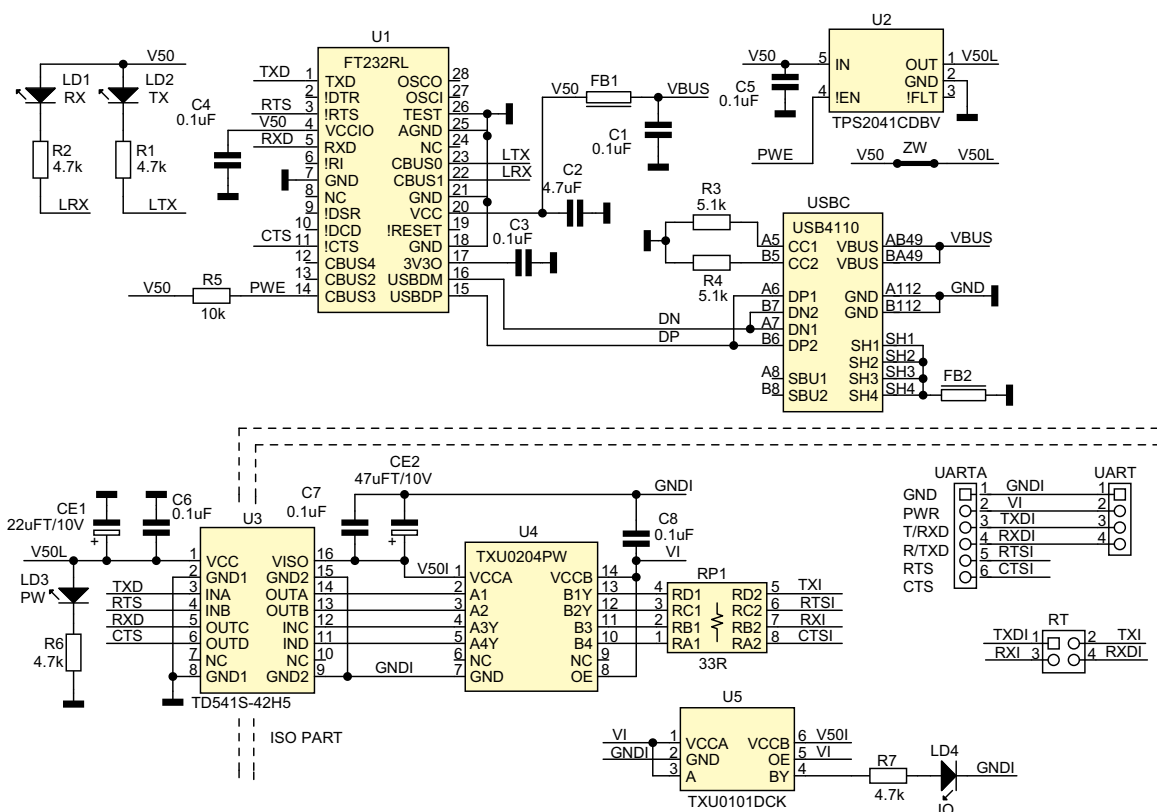
Izolowany konwerter USB-UART z przetwornicą poziomów

Konwerter USB/UART jest narzędziem powszechnie stosowanym podczas uruchamiania oprogramowania lub zarządzania systemami operacyjnymi przez SSH. Opisany konwerter zawiera ponadto kilka dodatkowych bloków funkcjonalnych, takich jak izolacja galwaniczna czy translator poziomów. Ten ostatni umożliwia używanie konwertera z napięciami 1,2...5 V, niedostępnymi w standardowych wersjach tego typu urządzeń, co zdecydowanie ułatwia prace uruchomieniowe.

Schemat konwertera pokazano na **rysunku 1**. Za konwersję USB/UART odpowiada popularny układ FT232RL, jego aplikacja nie odbiega od zalecanej w notcie katalogowej. Interfejs USB doprowadzony



jest do gniazda USB-C pracującego ze zmniejszoną liczbą wyprowadzeń, zapewniającą zgodność z USB2.0. Dioda LD3 sygnalizuje obecność zasilania USB-C, a diody LD1 i LD2 – obecność transmisji. Układ U2 to przełącznik zasilania typu TPS2041, sterowany sygnałem PWEN – kluczuje on napięcie V50L, doprowadzone do modułu



Rysunek 1. Schemat ideowy konwertera

Wykaz elementów:

Rezystory:

R1, R2, R6, R7: 4,7 kΩ (SMD 0603, 1%)
R3, R4: 5,1 kΩ (SMD 0603, 1%)
R5: 10 kΩ (SMD 0603, 1%)
RP1: drabinka 4×33 Ω 1% (CRA06S08)

Kondensatory:

C1, C3...C8: 100 nF (SMD 0603, X7R, 10 V)
C2: 4,7 μF (SMD 0603, X7R, 10 V)
CE1: 22 μFT/10V tantalowy (SMD A/3216)
CE2: 47 μFT/10V tantalowy (SMD B/3528)

Półprzewodniki:

LD1: dioda LED żółta (SMD 0603)
LD2: dioda LED czerwona (SMD 0603)
LD3: dioda LED zielona (SMD 0603)
LD4: dioda LED niebieska (SMD 0603)
U1: FT232RL (SSOP28)
U2: TP52041CDBV (SOT-23-5)
U3: TD541S-42H5 (DFN16)
U4: TXU0204PW (TSSOP14, 065)
U5: TXU0101DCK (SC70-6)

Pozostałe:

FB1, FB2: dławik ferrytowy (SMD 0603, typ BLM18PG121SZ1D)
RT: złącze IDC4 + 2 zwory
UARTA: złącze SIP6 (2,54 mm)
USBC: gniazdo USB-C (USB4110, GTC)

izolatora cyfrowego U3 typu TD541S-42H5, aby zapewnić minimalny pobór prądu podczas inicjalizacji konwertera. U3 to nowe opracowanie firmy Mornsun, integrujące w jednej obudowie izolowaną przetwornicę napięcia oraz cztery kanały izolatorów cyfrowych. W ramach rodziny TD541S-4x dostępne są izolatory o czterech kanałach transmisji, ale o zaleźnym od modelu, ustalonym kierunku transmisji DI->DO. Dostępne są konwertery w wersjach: 4:0, 3:1, 2:2, 1:3 i 0:4 kanałów DI:DO, co kompleksowo rozwiązuje separację interfejsów GPIO i SPI, niezależnie od strony interfejsu, z której zasilany jest układ. TD541S-42H5 zawiera po dwa izolatory pracujące w każdym z kierunków (2:2), co umożliwia realizację interfejsu UART z potwierdzeniem sprzętowym. Wybór wersji zasilanej napięciem 5 V upraszcza obwód zasilania części pierwotnej konwertera, gdyż FT232RL pracuje w tym standardzie napięciowym bez dodatkowych elementów zewnętrznych (VCCIO=5 V).

Po izolacji galwanicznej sygnały UART doprowadzone są do translatora poziomów U4 typu TXU0204. Zastosowanie dodatkowego translatora podyktowane jest koniecznością zapewnienia poprawnej pracy konwertera zarówno z układami niskonapięciowymi 1,2/1,8 V, jak i tymi zasilanymi napięciem 2,5...5 V – coraz więcej nowoczesnych układów SoC lub komputerów jednopłytkowych pracuje bowiem ze standardem napięciowym GPIO 1,2 V lub 1,8 V, podczas gdy dostępne na rynku czujniki i moduły rozszerzeń często wymagają 3,3 V lub nawet 5 V.

Dodatkowym wymogiem przy projektowaniu konwertera był minimalny prąd pobierany z układu uruchamianego, aby nie zakłócać zasilania układów o niskim poborze mocy. Strona A konwertera U4,

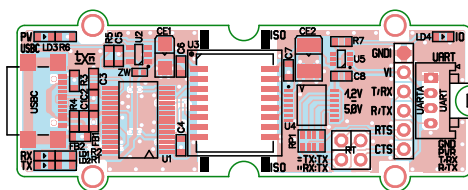
połączona z izolatorem U3, zasilana jest napięciem 5 V pochodzącym z wbudowanej w U3 przetwornicy, zaś strona B – z napięcia VI pochodzącego z uruchamianego układu. Układ U5 steruje diodą LD4 sygnalizującą obecność zasilania z układu uruchamianego. Zastosowanie dodatkowego konwertera podyktowane jest koniecznością minimalizacji poboru prądu z tegoż układu oraz zapewnienia napięcia koniecznego do zaświecenia LED, co przy współpracy z układami zasilanymi napięciem 1,2/1,8 V okazuje się niewykonalne. W modelu LD4 zasilana jest napięciem VI=5 V. Całkowity, maksymalny, statyczny pobór prądu z układu uruchamianego za pomocą konwertera nie przekracza 20 μA. W rzeczywistości jego natężenie jest znacznie niższe i przy prędkości 921600 bps i zastosowaniu transmisji z potwierdzeniem (zwarte RXD/TXD, CTS/RTS) pobór prądu nie przekraczał 2 μA przy zasilaniu 1,2 V i 6 μA przy zasilaniu 5 V. Układy TXU0204 i TXU0101 gwarantują przejście wyprowadzeń IO w stan wysokiej impedancji, jeżeli jedno z napięć zasilających spadnie poniżej 100 mV – gwarantuje to brak przepływów wstecznych GPIO lub niepożądanego efektu zasilania zewnętrznego urządzenia przez konwerter, co w klasycznych rozwiązaniach mogłoby stwarzać problemy np. z poprawnym restartem procesorów zasilanych pasożytniczo przez GPIO. Wszystkie sygnały UART konwertera wyprowadzono na złącze UARTA, a podstawowe sygnały RXD/TXD – dodatkowo także na złącze w standardzie Grove (UART). Zwora RT umożliwia zamianę wyprowadzeń RXD i TXD, co niweluje problemy z różnymi sposobami wyprowadzania ich na złącza.

Układ zmontowano na dwustronnej płytce drukowanej, rozmieszczenie elementów zaprezentowano na **rysunku 2**. Montaż nie wymaga szczegółowego opisu.

Zmontowany moduł konwertera pokazano na **fotografii tytułowej**.

Układ po podłączeniu do komputera z systemem Windows wykrywany jest automatycznie i nie wymaga dodatkowej konfiguracji. Sprawdzenie działania konwertera można wykonać przy pomocy terminala portu szeregowego, weryfikując poprawność transmisji po zwarcie wyprowadzeń RXD/TXD i doprowadzeniu napięcia zasilania 1,2...5 V do wyprowadzeń VI/GND.

Adam Tatuś, EP



Rysunek 2. Rozmieszczenie elementów na PCB modułu

Silniki krokowe FAULHABER

Wyzwania? Zapraszamy do Środka!

Wysoka wydajność silnika
krokowego serii DM 66200H
otwiera nowe możliwości integracji
w wymagających zastosowaniach.

faulhaber.com/ringstepper/en



WE CREATE MOTION



W ofercie AVT*

AVT6088

Najważniejsze parametry:

- zakresy napięcia i prądu wejściowego ustalone przez dobór elementów (domyślnie: 4...20 V, 0...1 A),
- typowa dokładność: 3%,
- napięcie zasilania układu: 3...5,5 V (typ. 5 V),
- pobór prądu: 2 mA (typ.),
- kalibracja skali przetwarzania za pomocą potencjometru,
- wyjście analogowe w standardzie Grove.

* **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.

Projekty pokrewne na stronie www.ep.com.pl

(aktywne linki do artykułów):

- Dwukanałowy multiplexer magistrali I²C zgodny z systemem Grove
- Cyfrowy termometr/termostat I²C zgodny z Grove
- Translator poziomów I²C Grove
- Ośmiokanałowy mostek master I²C/1-Wire Grove
- Moduł precyzyjnego zegara RTC z interfejsem Grove

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

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

Moduł przetwornika mocy DC w standardzie Grove

Prezentowany układ to miniaturowy moduł przetwornika mocy prądu stałego, przydatny w robotyce amatorskiej lub w układach automatyki domowej do kontroli mocy pobieranej przez odbiorniki, co jest szczególnie istotne w przypadku zasilania bateryjnego lub akumulatorowego.

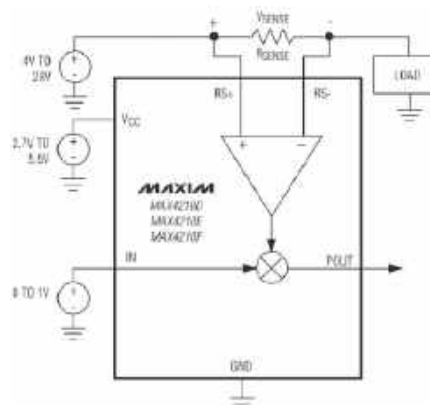
Minimoduł oparto na scalonym przetworniku mocy typu MAX4210(x) firmy Analog Devices. Strukturę wewnętrzną układu pokazano na **rysunku 1**. Układ mierzy moc, mnożąc wartość napięcia pomiarowego (pochodzącego z zewnętrznego dzielnika rezystorowego, który należy połączyć z wejściem IN) przez sygnał prądu, mierzony pośrednio jako spadek napięcia na rezystorze pomiarowym (boczniku) pomiędzy końcówkami RS+ i RS-. Zmierzone napięcia wymnażane są wewnątrz układu i, po przeskalowaniu o wartość zależną od podtypu układu MAX4210D/E, wyprowadzane na wyjście Pout. Gotowa, przeliczona wartość mocy dostępna jest na wyjściu Pout układu, co uwalnia nas od konieczności osobnego pomiaru napięcia i prądu oraz dodatkowych kalkulacji. W ten sposób moduł pozwala oszczędzić

jedno wejście analogowe mikrokontrolera, a dodatkowo uwalnia współpracujący procesor od konieczności wykonywania dodatkowych obliczeń.

Schemat ideowy modułu pokazano na **rysunku 2**. Do złącza IN doprowadzone jest napięcie zasilające obciążenie, a do złącza OUT – monitorowane obciążenie. Obsługiwany zakres napięć zasilania wynosi 4...28 V, co pozwala na pomiar w układach zasilanych bateryjnie lub akumulatorowo. Rezystor RS dobierany jest w zależności od wymaganego prądu obciążenia. Maksymalna różnica napięć RS+/RS- w przypadku układów MAX4210D/E wynosi 150 mV. Do wejścia IN, poprzez dzielnik R1...R4, doprowadzone jest napięcie zasilania, przeskalowane do zakresu 0...1 V. W modelu dzielnik 1:20 ustala maksymalne napięcie zasilania na wartość 20 V, a rezystor



RS1=0,1 Ω definiuje zakres prądu pomiarowego równy 0...1 A, co umożliwi pomiar mocy w zakresie ~0,5...20 W. Napięcie wyjściowe, odpowiadające przetworzonej



Rysunek 1. Struktura wewnętrzna MAX4210(x) za notą Analog Devices

Wykaz elementów:**Rezystory:**

- R1: 100 kΩ (SMD 0603, 1%)
- R2, R3: 200 kΩ (SMD 0603, 1%)
- R4: 1,5 MΩ (SMD 0603, 1%)
- R5: 22 kΩ (SMD 0603, 1%)
- R6: 10 kΩ (SMD 0603, 1%)
- R7: 10 Ω (SMD 0603, 1%)
- RS1: 0,1 Ω (SMD 2512, 1 W, 1%)
- RV1: potencjometr wielobrotowy 10 kΩ (VR-64W)

Kondensatory: (SMD 0603, X7R, 10 V)

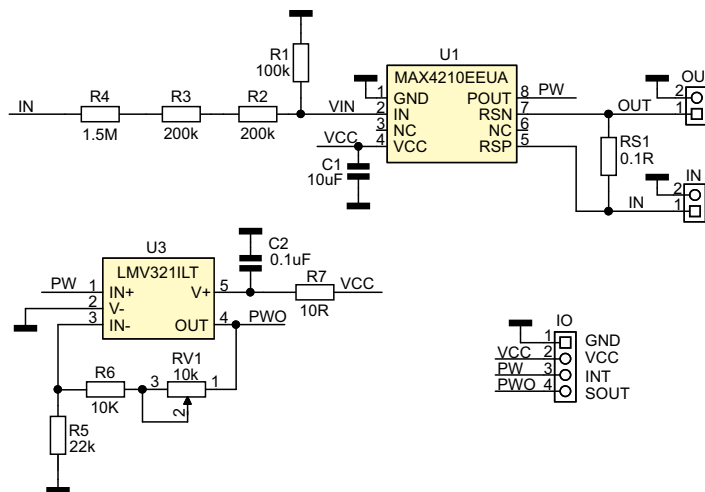
- C1: 10 μF
- C2: 100 nF

Półprzewodniki:

- U1: MAX4210EEUA (UMAX8)
- U3: LMV321ILT (SOT-23-5)

Pozostałe:

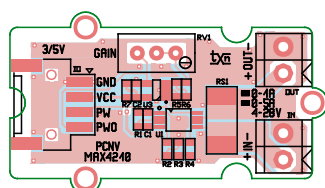
- IO: złącze kątowe Grove (SMD, typ HY2.0-4P)
- IN, OUT: złącze śrubowe 2 pin. (3,5 mm, typ DG381-3.5-2)



Rysunek 2. Schemat przetwornika mocy

Tabela 1. Wyniki pomiarów prototypowego egzemplarza modułu

Napięcie wejściowe Uwe [V]	Prąd wejściowy Iwe [A]	Moc wejściowa obliczona Pwe [W]	Napięcie wyjściowe obliczone Upo [V]	Napięcie wyjściowe zmierzone Upo [V]	Napięcie wyjściowe zmierzone Upwo [V]	Błąd napięcia wyjściowego dp Upo [%]	Moc wyjściowa (PWO) Ppwo [W]	Błąd (PWO) dp Ppwo [%]
20,0	1,00	20,00	2,500	2,532	4,00	1,28	20,00	0,0
15,0	1,00	15,00	1,875	1,892	2,99	0,91	14,95	-0,3
10,0	1,00	10,00	1,250	1,253	1,98	0,24	9,90	-1,0
5,0	1,00	5,00	0,625	0,613	0,97	-1,92	4,86	-3,0
20,0	0,50	10,00	1,250	1,270	2,01	1,60	10,05	0,5
15,0	0,50	7,50	0,938	0,950	1,50	1,33	7,50	0,0
10,0	0,50	5,00	0,625	0,631	0,99	0,96	4,95	-1,0
5,0	0,50	2,50	0,313	0,311	0,49	-0,48	2,45	-2,0
20,0	0,10	2,00	0,250	0,251	0,40	0,40	2,00	0,0
15,0	0,10	1,50	0,188	0,187	0,30	-0,27	1,50	0,0
10,0	0,10	1,00	0,125	0,124	0,20	-0,80	1,00	0,0
5,0	0,10	0,50	0,063	0,061	0,10	-2,40	0,50	0,0



Rysunek 3. Rozmieszczenie elementów na płytce modułu

mocy, dostępne jest na wyprowadzeniu POUT. Stopień wzmocnienia zbudowany w oparciu o wzmacniacz U3 (o wzmocnieniu regulowanym potencjometrem RV1) umożliwia dodatkowe przeskalowanie napięcia wyjściowego PWO. Wartość napięcia wyjściowego określona jest wzorem:

$$V_{\text{pout}} = A_{\text{vpout}} \cdot V_{\text{sens}} \cdot V_{\text{in}}$$

Wartość A_{vpout} jest ustalona w zależności od podtypu układu: w przypadku MAX4210D wynosi 16,67, a w przypadku MAX4210E: 25. Po uwzględnieniu wartości rezystorów zastosowanych

w modelu otrzymujemy: $V_{\text{pout}} = 0,125 \cdot P_{\text{we}}$. Optymalizując układ pod kątem własnych zastosowań należy zwrócić uwagę, aby nie przekraczać zakresu napięć wejściowych układu mnożącego oraz pamiętać o ograniczeniu napięcia wyjściowego układu wzmacniacza, związanym z wartością napięcia zasilania. Układ wymaga zasilania 3...5,5 V, pobór prądu wynosi ok. 2 mA (bez obciążenia). Napięcie zasilania oraz napięcia wyjściowe PW, PWO wyprowadzone są na złącze IO zgodne z Grove.

Moduł zmontowano na dwustronnej płytce drukowanej. Rozmieszczenie elementów pokazano na **rysunku 3**. Sposób montażu jest klasyczny i nie wymaga opisu. Zmontowany moduł można zobaczyć na **fotografii tytułowej**.

Uruchomienie przetwornika wymaga ustalenia współczynnika przetwarzania na wyjściu PWO potencjometrem RV1. Dokładność sygnału na linii PW jest

określona przez tolerancje rezystorów R1...4 i ustawienie RS1. Po podłączeniu napięcia zasilającego obciążenie do złącza IN oraz samego obciążenia (na czas uruchomienia najlepiej użyć sztucznego obciążenia o regulowanym prądzie) do złącza OUT trzeba jeszcze doprowadzić zasilanie modułu (5 V) do pinu VCC. Następnie dla kilku różnych wartości mocy wejściowej sprawdzamy działanie układu. Ustalając np. napięcie 20 V i prąd 1000 mA (20 W) regulujemy potencjometrem RV1 napięcie na wyjściu PWO tak, aby wynosiło ono dokładnie 4,00 V, co daje współczynnik przetwarzania 5 [W/V]. Przykładowe wartości uzyskane podczas pomiaru modelu pokazano w **tabeli 1**.

Wartości błędów nie przekraczają 3%, co można uznać za wartość wystarczającą dla typowych, technicznych pomiarów mocy.

Adam Tatuś, EP

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>

Arduino R4 Nano

– nie jest źle, a nawet lepiej

W połowie wakacji zespół Arduino wprowadził na rynek nową płytkę w formacie Arduino Nano, tym razem z 32-bitowym procesorem z rodziny Renesas RA4M1. Po płytkach z ESP32, RP2040, MG24 mamy zatem do dyspozycji kolejny niewielki moduł z 32-bitowym procesorem.

Podobnie jak w przypadku większych modeli Arduino R4 Minima i Arduino Uno R3, płytka Nano R4 ma być alternatywą dla leciwego już Arduino Nano. Dobrą wiadomością jest fakt, że zachowano dokładnie ten sam kontroler R7FA4M1AB3CFM, co w większej wersji. Wróży to zachowanie zgodności z wcześniejszymi projektami, które można będzie przenieść na mniejszą (pod względem rozmiarów) platformę. Dla porównania w tabeli 1 zestawiono podstawowe parametry Uno R4 Minima, Nano R4 i Nano.

Nano R4 bazuje na układzie R7FA4M1AB3CFM w obudowie LQFP-64, którego budowę wewnętrzną pokazano na rysunku 1. Konkurencyjne płytki XIAO RA4M1 są budowane w oparciu o R7FA4M1AB3CNE w obudowie QFN-48, o zmniejszonej liczbie wyprowadzeń GPIO, co wymaga dodatkowej konfiguracji Arduino i pewnych zmian w oprogramowaniu – ale w zamian otrzymujemy tanią (koszt to około 6 €) i naprawdę niedużą płytkę. Nano R4 jest

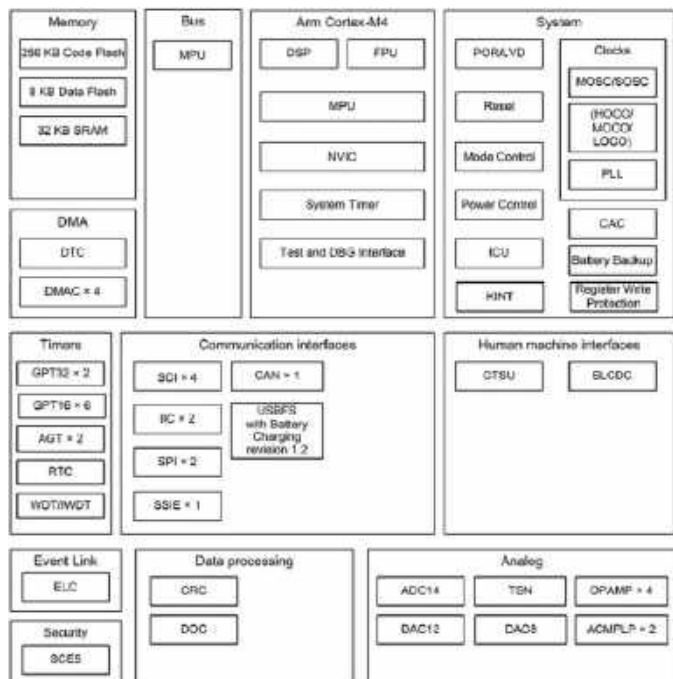


dostępne w wersjach z wlutowanymi złączami goldpin oraz bez nich (są natomiast dołączone do płytki w opakowaniu), a różnica w cenie to trochę ponad 1 €. Tak jak wcześniej opracowane wersje z ESP32 czy RP2040, płytka – oprócz wlutowania złączy – umożliwia bezpośredni montaż SMT na płycie bazowej urządzenia, gdyż wyposażona jest w metalizowane złożone pady na krawędziach (tzw. castellated holes – przyp. red.).

Z racji swojej przemysłowej i motoryzacyjnej specjalizacji, firma Renesas oferuje procesor w rozszerzonym zakresie temperatur, co zawsze gwarantuje większą stabilność rozwiązania (o ile inne komponenty R4 wytrzymają $-40...85^{\circ}\text{C}$). W porównaniu do R4 Minima, pomimo mniejszej płytki otrzymujemy większą liczbę wyprowadzonych linii GPIO, co jest zawsze mile widziane w projektach – szczególnie

Tabela 1. Porównanie wyposażenia płytek Uno R4 Minima, Nano R4 i Nano

Płytki		Arduino Uno R4 Minima	Arduino Nano R4	Arduino Nano
Nr katalogowy		ABX00080	ABX00143	A000066
Mikrokontroler		Renesas RA4M1 Arm Cortex-M4 R7FA4M1AB3CFM	Renesas RA4M1 Arm Cortex-M4 R7FA4M1AB3CFM	ATmega328P
USB/programator		USB-C/złącze SWD	USB-C/SWD pogo	USB-B
GPIO	DIO	14	22	20
	ADC	6 (14 bit.)	8 (14 bit.)	8 (10 bit.)
	DAC	1 (12 bit.)	1 (12 bit.)	0
	PWM	6	6	6
	Maksymalny prąd GPIO	8 mA	8 mA	20 mA
Porty komunikacyjne	UART	1	1	1
	I ² C	1	2 (złącze Qwiic)	1
	SPI	1	1	1
	CAN	1	1	0
Zasilanie	VCC	5 V	5 V	5 V
	VIN	6...24 V	6...21 V	7...12 V
Taktowanie	CPU CLK	48 MHz	48 MHz	16 MHz
Pamięć	Flash	256 kB	256 kB	32 kB
	RAM	32 kB	32 kB	2 kB
	EEPROM	Data Flash 8 kB	Data Flash 8 kB	1 kB
Cena		18 €	14 €	27 €
Pozostałe		wzmacniacz operacyjny, RTC	wzmacniacz operacyjny, RTC, wbudowany kwarc 16 MHz, dioda RGB, konwerter poziomów 3,3 V dla I ² C0 wyprowadzonego na złącze QWIIIC	



Rysunek 1. Budowa procesora z rodziny R4M1 (za notą Renesas)

że przybyło wyprowadzeń analogowych, pojawił się dodatkowy port I²C oraz magistrala CAN. W porównaniu z Nano zmieniono złącze USB micro na USB-C, co jest raczej rzeczą oczywistą i przystaje do promowanych na całym świecie rozwiązań. Oprócz diody użytkownika płytkę wyposażoną jest w diodę LED RGB – na szczęście żadna z diod nie blokuje pinów GPIO wyprowadzonych na złącza, są one bowiem podłączone do nieużywanych wyprowadzeń sporej obudowy TQFP-64 (a i tak zostało jeszcze 14 wolnych linii).

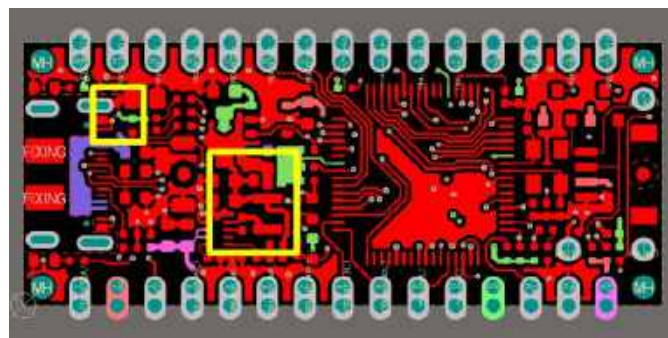
Sporą zmianą jest wyposażenie procesora w taktowanie zewnętrznym kwarcem 16 MHz – w przeciwieństwie do Uno R4, w którym zastosowano wbudowany generator. Jeżeli rozwiązanie to zostanie dobrze zaimplementowane, zdecydowanie zwiększy ono dokładność operacji krytycznych pod względem czasu wykonywania kodu. Zegar czasu rzeczywistego ma możliwość podtrzymania przy pomocy zewnętrznej baterii 3 V. Nie przewidziano jednak na płytce ani uchwytu baterijnego, ani pełnoprawnego złącza do podłączenia takiego źródła energii, wyprowadzenie dodatnie musi więc być przylutowane bezpośrednio do przewidzianego padu na płytce, podobnie wyprowadzenie ujemne trzeba na własną rękę dolutować do masy. To niezbyt eleganckie rozwiązanie – są przecież dostępne małe złącza JST1.0, takie jak w Raspberry Pi 5. Jeżeli chodzi o zgodność wyprowadzeń z Uno, to jest ona praktycznie zachowana, z wyjątkiem jednego wyprowadzenia BOOT, które w Uno pełni funkcję RST (reset).

W przypadku układu zasilania poprawiono zakres napięć, przy których poprawnie pracuje wbudowana przetwornica. W Nano

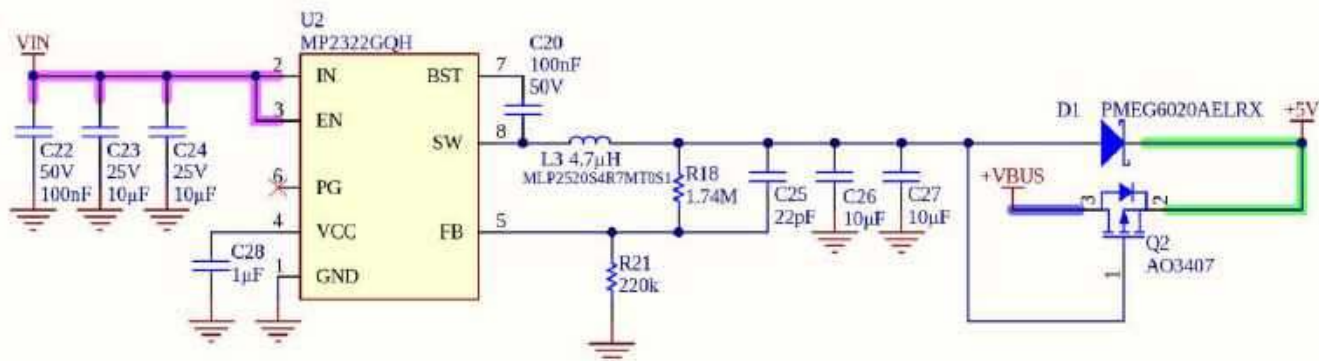
R4 sięga on 21 V, choć sama przetwornica może pracować w jeszcze szerszym zakresie. Rozwiązano też problem oscylacji przetwornicy i zaniżonego napięcia zasilania +5 V, co zrealizowano poprzez zmianę sposobu kluczowania napięć pochodzących z zasilacza (dioda D1) i portu USB (tranzystor Q2) – patrz **rysunek 2**. Zmiany dotyczą też obwodu napięcia 3,3 V, który w Minima zasilany był z wbudowanego w procesor RA4M1 stabilizatora LDO, co niepotrzebnie nagrzewało strukturę i w przypadku zwarcia mogło okazać się niebezpieczne dla procesora. W Nano R4 przewidziano osobny stabilizator 3,3 V (U3 typu AP2112-3.3), co ciekawe – kluczowany za pomocą wyprowadzenia P500 procesora. To znaczące zmiany na lepsze, eliminujące błędy projektowe z Minima.

Jedynie, do czego można się przyczepić, to fakt, że projektanci Arduino w dalszym ciągu nie wykorzystują szerokiego zakresu zasilania procesorów R4M1 (1,6...5,5 V) i „na sztywno” zasilają je napięciem 5 V. Z Nano R4 znikło złącze programatora SWD – jest teraz dostępne tylko w formie padów na dolnej stronie płytki (tylko jak do niego się dostać, gdy płytka będzie wlutowana?). Z powodu kilku pinów trzeba będzie zatem „piec” płytkę na paście, zamiast szybko i wygodnie polutować metalizowane pady krawędziowe zwykłą lutownicą.

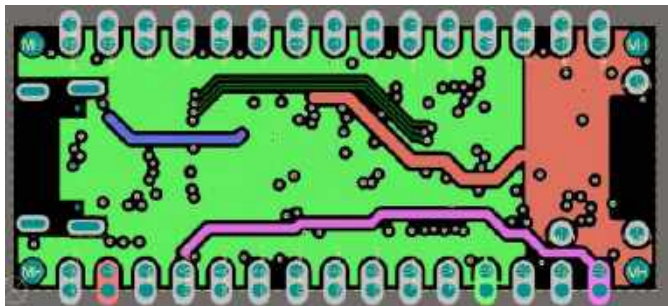
Płytkę wspiera tryb USB-HID, umożliwiający emulację klawiatury i myszy, co jest dużą zaletą Nano R4. Jeżeli chodzi o projekt PCB, został nieco poprawiony. Nano R4 wykonano w oparciu o 4-warstwową płytkę drukowaną, z wydzielonymi płaszczyznami masy i zasilania. W dalszym ciągu całe zasilanie USB (VBUS) idzie poprzez jedną, „biedną” przelotkę o średnicy 8 milsów, co jest dosyć ryzykowne (**rysunek 3**). Podobnie rachityczne są połączenia głównego obwodu kluczowania przetwornicy 5 V. Dostatecznie wykorzystano według mnie także warstwę zasilania, gdzie zamiast użycia wylewek masy (GND) o małej impedancji i dobrym, rozproszonym filtrowaniu zdecydowano się na zastosowanie wąskich i długich ścieżek, zajmując większą część warstwy wylewki GND, co pokazano na **rysunku 4**. Widoczne są też mikroprzerwy w wylewce masy w okolicach przelotek. Ostatnim smaczkiem jest magistrala USB, gdzie pady testowe tworzą ładne „antenki”



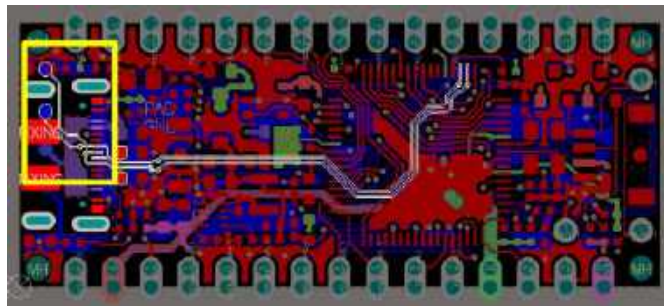
Rysunek 3. Warstwa TOP z jedną, mizerną przelotką zasilania VBUS i wąskimi ścieżkami obwodu kluczowania przetwornicy



Rysunek 2. Obwody kluczowania napięć zasilania (fragment oryginalnego schematu z dokumentacji firmy Arduino)



Rysunek 4. Warstwa zasilania



Rysunek 5. „Antenki” na magistrali USB

na sygnałach DP+ i DP-, a wystarczyło je przenieść w inne miejsce (rysunek 5). W Mininie było serduszko – tutaj są czułka, tylko dławczego na całkiem szybkiej magistrali? Parafrazując powiedzenie „nie próbuj tego w domu” możemy zatem powiedzieć „nie naśladowanie projektu tych płytek”. Zaleta wersji Nano R4 to natomiast znacznie skrócone połączenia wbudowanego, 14-bitowego przetwornika A/D – chociaż bliskość dławika przetwornicy i ferrytów filtrujących zasilanie części analogowej nie napawa optymizmem i zapewne trudno byłoby wykorzystać pełną, 14-bitową rozdzielczość konwertera. Płytkę pozbawioną jest jakichkolwiek opisów na warstwie TOP, co nieco utrudnia jej sprawne łączenie na płytkach stykowych.

Nano R4 wyposażono w dwa interfejsy komunikacyjne UART: pierwszy – wirtualny, realizowany przez USB, drugi – sprzętowy, dostępny na wyprowadzeniach RX/TX. Magistrala I²C, doprowadzona do złącza SIP, jest zgodna z poziomami logicznymi 5 V. Należy zwrócić uwagę, że płytkę nie zawiera rezystorów podciągających sygnały SDA i SCL, wyprowadzone na złącza Analog A4/5. Drugi interfejs I²C – przez konwerter poziomów 5 V/3,3 V – wyprowadzono wraz z zasilaniem 3,3 V na złącze QWIIIC. To szczególnie istotne, gdyż proponowane przez zespół Arduino moduły rozszerzeń MODULINO, są właśnie w takie złącze wyposażone.

Moduł obsługuje ponadto interfejs CAN, dostępny na wyprowadzeniach D5 (CANRX), D4 (CANTX) i wymagający tylko dodania zewnętrznego drivera magistrali (np. SN65HVD230) oraz zastosowania biblioteki Arduino_CAN.h. Interfejsy szeregowo uzupełnia szyna SPI, dostępna na złączach krawędziowych (jak w Nano). Niestety w dalszym ciągu niestabilne – tak jak w przypadku Uno R4 Minima – jest ładowanie aplikacji. Pomimo widoczności w Menedżerze Urządzeń, co pewien czas zaprogramowanie płytki z poziomu środowiska Arduino (ver. 2.3.6) jest niemożliwe. Tymczasowo problem rozwiązuje uruchomienie loadera DFU i restart środowiska, a czasem całego systemu operacyjnego.

Na zakończenie z ciekawości porównałem wydajność płytek rodziny Nano oraz mniejszych rozmiarowo Xiao, przy pomocy szkicu testowego obliczającego liczbę Pi przy użyciu biblioteki *MonteCarloPi by cygig v0.8.3*. Przykłady dla każdego procesora zostały skompilowane w trybie używającym tylko jednego rdzenia (w przypadku procesorów wyposażonych w kilka rdzeni). Otrzymane wyniki zebrano w tabeli 2.

Tak jak było do przewidzenia, R4 Nano ma wydajność obliczeniową znacznie większą niż płytki Nano z procesorami 328. Niestety jest ona znacznie niższa od osiągnięć RP2040 z Raspberry Pico, które jest najszybszym rdzeniem obliczeniowym, wyprzedzającym nawet procesory z rodziny ESP.

Podsumowanie

Płytkę Nano R4 nie jest idealna, ale jej projekt niweluje kilka irytujących błędów z większej wersji Uno R4, zachowując zgodność z klasycznym Nano. Zapewne więc znajdzie swoją grupę odbiorców. Nie licząc ponownie występującej tu niedogodności (a może – wpadki?) z podtrzymaniem baterijnym RTC, płytkę po kilku dniach testów wydaje się być godna polecenia. Jeżeli ktoś jeszcze nie używał procesora R4, wydajność Nano już mu nie wystarczy, a dodatkowo nie jest ograniczony wymogiem zachowania zgodności z formatem Uno – wersja Arduino Nano R4 okaże się dlań sensowniejszym wyborem od Minima (nie wspominając o kuriozalnym modelu z Wi-Fi). I to nie tylko ze względu na nieco lepsze wyposażenie i sensowniejszy projekt, ale też znacznie niższą cenę (13,4 €/22 €). Jeżeli jednak zależy nam na znacznie większej wydajności, należy kierować się w stronę nowocześniejszych, wielordzeniowych procesorów ESP lub RP2040 – tym bardziej, że dostępne zestawy są nawet o połowę tańsze.

Adam Tatuś, EP

Tabela 2. Zestawienie wydajności procesorów

	Szacowana wartość liczby Pi (500000 pętli)	Czas obliczeń [ms]	Szacowana wartość liczby Pi do 5 miejsc po przecinku	Czas obliczeń [ms]	Liczba pętli
R4 Nano	3,141464	21091	3,141592	1415	20559
Nano	3,140336	101381	3,141593	2491	9040
Nano Every 328	3,140336	99424	3,141593	2457	9040
Nano Every 4809	3,140336	99424	3,141593	2457	9040
Nano DFROBOT 328	3,140336	101381	3,141593	2491	9040
Arduino RP2040 Connect	3,141464	3175	3,141592	229	20559
Nano 33 BLE	3,141464	14688	3,141592	1014	20559
Nano 33 IOT	3,141464	39115	3,141592	2798	20559
XIAO ESP32S3	3,141464	3503	3,141592	235	20559
XIAO ESP32C3	3,141464	5441	3,141592	394	20559
XIAO ESP32C6	3,141464	3863	3,141592	274	20559
XIAO MG24	3,141464	23746	3,141592	1576	20559
Arduino Micro	3,140336	101907	3,141593	2503	9040

Druk 3D – moda czy niezbędne narzędzie wspierające w sytuacjach awaryjnych?

Nieoczekiwane uszkodzenie elementu i zatrzymanie linii produkcyjnej to scenariusz, w którym liczy się szybkość reakcji i trafne decyzje. Awaria podzespołów czy naturalne zużycie materiałów mogą w każdej chwili sparaliżować pracę zakładu. Regularna kontrola i systemy wczesnego wykrywania problemów pomagają zminimalizować to ryzyko, lecz nie są w stanie go całkowicie wykluczyć. Dlatego w krytycznych momentach coraz większą rolę odgrywa druk 3D – umożliwia on błyskawiczne przygotowanie części zamiennych i sprawne przywrócenie maszyn do działania.

Cichy bohater w branży przemysłowej

„Druk 3D to już nie chwilowa moda, lecz sprawdzone narzędzie wykorzystywane w przemyśle do zwiększania elastyczności i wydajności produkcji”, mówi Florian Ebner, wieloletni ekspert ds. druku 3D w Conrad Electronic. Jego zdaniem wiele firm wciąż nie dostrzega pełnego potencjału tej technologii. „Gdy liczy się czas, a awaria maszyny zatrzymuje produkcję z powodu uszkodzonego elementu, możliwość lokalnego wydrukowania części zamiennych, np. uchwytu, koła zębatego czy zawiasu, staje się ogromnym atutem” – wyjaśnia. Zwraca jednak uwagę, że w przypadku samodzielnie produkowanych części należy uwzględnić kwestie związane z gwarancją oraz ewentualnym wykorzystaniem cudzej własności intelektualnej.

Druk 3D w aplikacjach przemysłowych – najważniejsze zalety

Profesjonalne drukarki 3D, zaprojektowane z myślą o zastosowaniach przemysłowych, oferują dużą elastyczność w produkcji skomplikowanych elementów – zarówno pod względem używanych



Fotografia 2. Awarie materiałów często prowadzą do nieplanowanego zapotrzebowania w procesie produkcji. Co zrobić, gdy brakuje materiałów zastępczych, a czas ucieka? (ilustracja wygenerowana przez ChatGPT-4o)

materiałów, jak i różnorodności form. „W parkach maszynowych wielu małych i średnich przedsiębiorstw wciąż pracują starsze modele urządzeń, do których nie są już dostępne części zamienne. W przypadku awarii idealnym rozwiązaniem okazuje się druk 3D – szybki i wygodny sposób na wykonanie brakującego elementu”, podkreśla ekspert firmy Conrad. Precyzyjne dopasowanie, krótki czas realizacji oraz możliwość produkcji niewielkich serii to kolejne atuty tej technologii, które są szczególnie istotne w sytuacjach problemów z dostępnością komponentów. Druk 3D może również wspierać efektywne zarządzanie częściami zamiennymi. Niektóre z nich nie muszą być już magazynowane w dużych ilościach – wystarczy je wytwarzać na bieżąco, co obniża koszty składowania i pozwala lepiej wykorzystać dostępny kapitał.

Inne zastosowania druku 3D

Optymalizacja zarządzania częściami zamiennymi to tylko jedno z wielu praktycznych zastosowań drukarek 3D w przemyśle. Dzięki specjalistycznym elementom konstrukcyjnym można modyfikować i indywidualnie dopasowywać narzędzia w taki sposób, aby umożliwić bezpieczny montaż bez ryzyka uszkodzeń. Druk 3D pozwala również na tworzenie specjalnych uchwytów i dopasowanych narzędzi montażowych, które znacząco skracają czas przebrojenia maszyn. Technologia addytywna wspiera także modernizację istniejących urządzeń i linii produkcyjnych (retrofit). Przykładowo zamiana łożysk ślizgowych na kulkowe staje się prostsza, podobnie jak wyposażenie starszych maszyn w nowoczesne czujniki, ponieważ drukarka 3D umożliwia szybkie wykonanie dedykowanych mocowań do nowych komponentów.

Więcej informacji: conrad.pl/zuzycie-materialu



Fotografia 1. Wyjątkowa elastyczność w produkcji części – drukarki 3D znajdują szerokie zastosowanie w środowiskach przemysłowych (ilustracja wygenerowana przez Midjourney v6.1)

Technologia Push-X: Nowa era łączenia przewodów

W dzisiejszym dynamicznie rozwijającym się świecie technologii, innowacje w dziedzinie połączeń elektrycznych stają się kluczowym elementem efektywności i niezawodności systemów. Jednym z najnowszych osiągnięć w tym sektorze jest technologia Push-X, która rewolucjonizuje sposób, w jaki łączymy przewody. W niniejszym artykule przyjrzymy się bliżej tej technologii, jej zaletom oraz zastosowaniom, a także zaprezentujemy najważniejsze informacje na temat serii XPC 1,5.

Czym jest technologia Push-X?

Technologia Push-X to nowoczesny system łączenia przewodów, który umożliwia szybkie i łatwe podłączanie zarówno sztywnych, jak i elastycznych przewodów, bez użycia narzędzi. Dzięki innowacyjnemu mechanizmowi użytkownicy mogą łączyć przewody z tulejką lub bez niej, co znacznie upraszcza proces instalacji i redukuje czas potrzebny na wykonanie połączeń. Łączenie przewodów odbywa się za pomocą specjalnie zaprojektowanej zapadki. W momencie kontaktu z odizolowaną końcówką przewodu, zapadka wyzwala zacisk, który mocuje przewód w komorze zaciskowej złącza, zapewniając trwałe i niezawodne połączenie mechaniczne oraz galwaniczne (fotografia 1).



Fotografia 1. Mechanizm mocowania przewodu Push-X w serii XPC

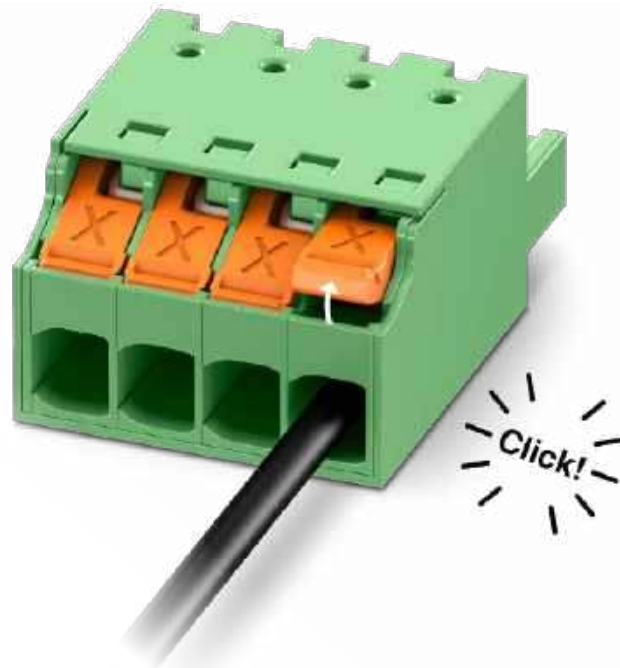
Kluczowe cechy technologii Push-X:

- **Bez narzędzi:** nowe złącza umożliwiają łatwe podłączanie przewodów bez konieczności używania dodatkowych narzędzi, co zwiększa wygodę użytkowania.
- **Wszechstronność:** technologia Push-X obsługuje różne typy przewodów, zarówno sztywne, jak i elastyczne, co czyni ją idealnym rozwiązaniem w różnych aplikacjach (fotografia 2).
- **Szybkość instalacji:** prosty w obsłudze mechanizm zmniejsza czas potrzebny na instalację, co jest szczególnie istotne w projektach wymagających szybkiego uruchomienia.

Zalety technologii Push-X

Technologia Push-X przynosi wiele korzyści zarówno dla producentów urządzeń i systemów, jak i dla użytkowników końcowych. Oto niektóre z nich:

- **Oszczędność czasu** – dzięki prostocie instalacji czas montażu jest znacznie krótszy w porównaniu do tradycyjnych metod łączenia. W praktyce zastosowanie technologii Push-X prowadzi do skrócenia czasu wykonywania instalacji nawet o 50% lub więcej.
- **Redukcja kosztów** – prostszy proces instalacji przekłada się na niższe koszty robocizny oraz mniejsze zużycie materiałów.



- **Zwiększona niezawodność** – systemy oparte na technologii Push-X charakteryzują się wysoką jakością połączeń, co zmniejsza ryzyko awarii.
- **Elastyczność zastosowań** – opisywana technologia znajduje zastosowanie w różnych branżach, takich jak automatyka przemysłowa, budownictwo czy elektronika.

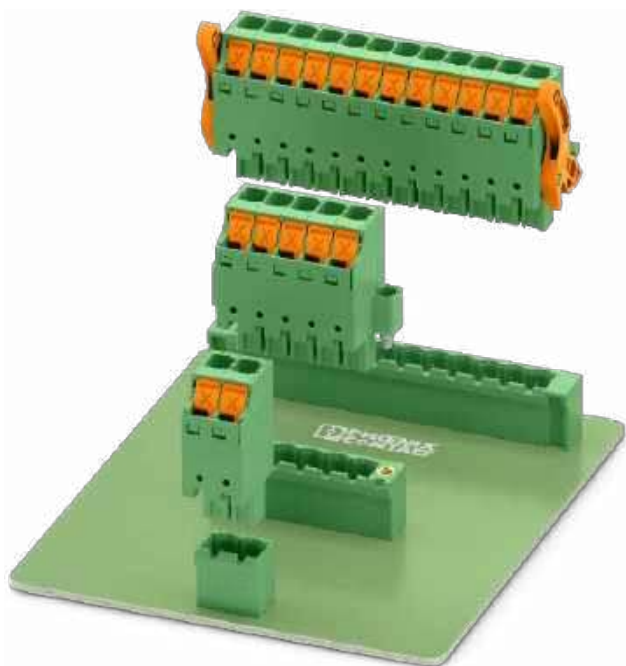
Zastosowania technologii Push-X

Push-X znajduje zastosowanie w wielu dziedzinach przemysłu. Oto kilka przykładów:

- **Automatyka przemysłowa** – szybkie podłączanie czujników i aktuatorów w systemach automatyki.
- **Budownictwo** – łączenie przewodów w instalacjach elektrycznych budynków.
- **Elektronika** – zastosowania w urządzeniach elektronicznych wymagających niezawodnych połączeń.



Fotografia 2. Technologia Push-X umożliwia beznarzędziowy montaż wszystkich typów przewodów



Fotografia 3. Wtyki serii XPC są kompatybilne ze standardowymi gniazdami COMBICON

Seria XPC 1,5

W ramach technologii Push-X dostępna jest seria XPC 1,5, która oferuje jeszcze więcej możliwości dla użytkowników. Seria ta została zaprojektowana z myślą o zapewnieniu wysokiej wydajności oraz niezawodności w różnych warunkach pracy. Obecnie dostępna jest wersja w rastrze 3,5 mm do przewodów o przekroju do 1,5 mm², a niebawem linia produktów zostanie wzbogacona o kolejne rastry i możliwość podłączenia przewodów o większych przekrojach. Wtyki są kompatybilne ze standardowymi gniazdami do PCB.

Kluczowe cechy serii XPC 1,5:

- **Wysoka jakość wykonania** – elementy serii XPC 1,5 są wykonane z materiałów odpornych na działanie wysokich temperatur oraz korozję.
- **Zgodność z normami** – produkty te spełniają rygorystyczne normy jakościowe i bezpieczeństwa, co czyni je odpowiednimi do zastosowań przemysłowych.
- **Łatwość montażu** – dzięki technologii Push-X montaż serii XPC 1,5 jest szybki i intuicyjny.

Technologia Push-X stanowi istotny krok naprzód w dziedzinie połączeń elektrycznych. Dzięki swojej prostocie i efektywności przynosi korzyści zarówno producentom, jak i użytkownikom końcowym. Seria XPC 1,5 to doskonały przykład zastosowania tej innowacyjnej technologii w praktyce. W miarę jak przemysł staje się coraz bardziej wymagający pod względem szybkości i niezawodności, technologie takie jak Push-X będą odgrywać kluczową rolę w przyszłości produkcji i automatyzacji.

Zamów próbki

Przekonaj się, jak szybkie i niezawodne może być beznarzędziowe łączenie przewodów z technologią Push-X: zamów bezpłatne próbki złączy XPC do PCB i przetestuj je w swojej aplikacji!

Zeskanuj kod QR

Dowiedz się więcej o technologii Push-X:
www.phoenixcontact.com/pl-pl/technologie-push-x



Piotr Gatuszka
 Device Connectors Design-in Engineer
 Phoenix Contact
pgaluszka@phoenixcontact.pl, tel. 882 082 744
www.strefaelektronika.com

REKLAMA



Idealne połączenia między płytkami

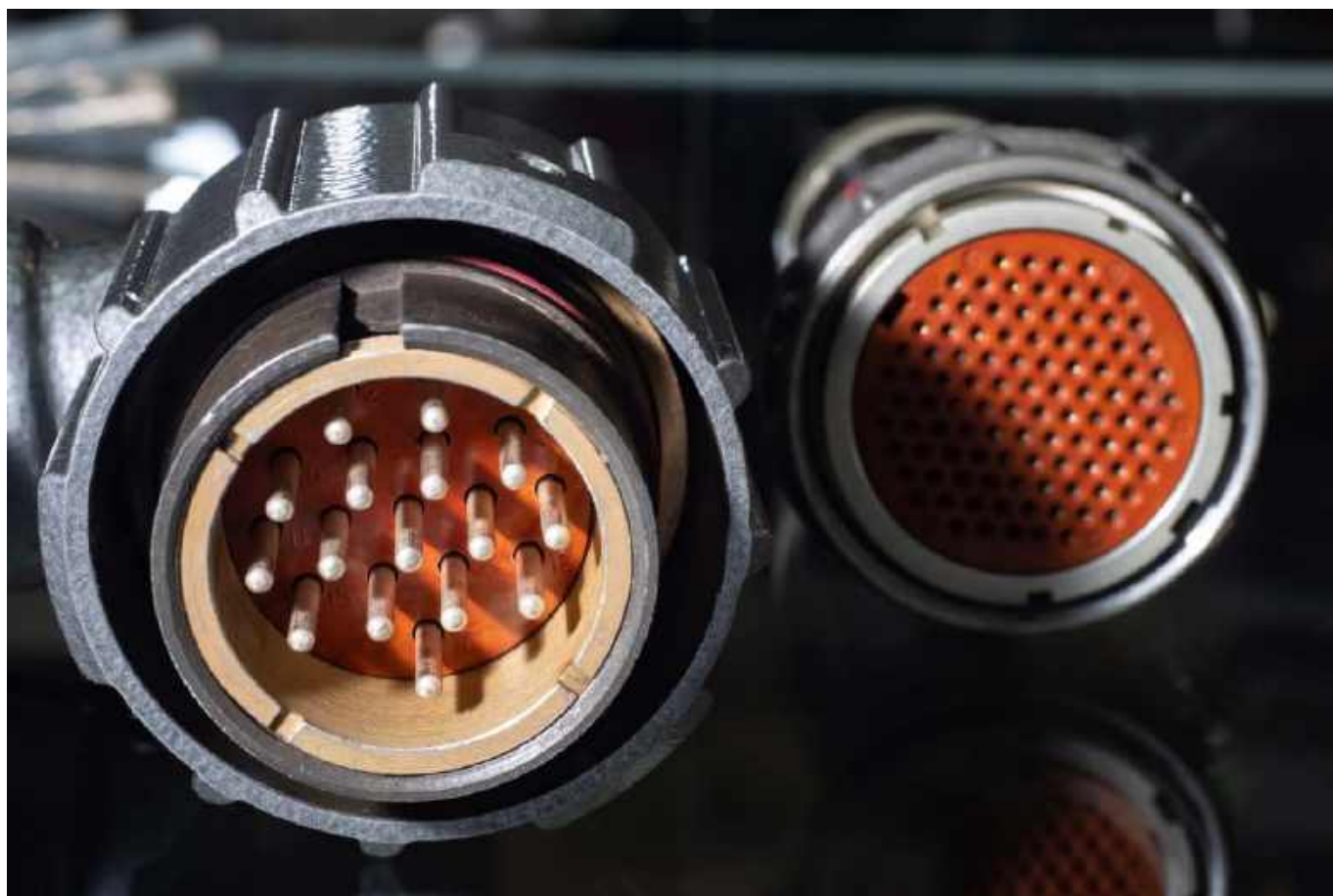
Niezawodne złącza serii FINEPITCH

- złącza w rastrach od 0,635 mm do 2,54 mm
- wykonanie ekranowane lub nieekranowane
- gwarancja bezawaryjnego łączenia płytek w różnych konfiguracjach orientacji w przestrzeni, odległości między laminatami lub ilości styków
- transmisja z dużą prędkością do 52 Gb/s



Odwiedź naszą stronę internetową i dowiedz się więcej!





Złącza do zastosowań specjalnych

– wysoka niezawodność w ekstremalnych środowiskach

Są takie projekty, w których ostatnim aspektem, na który zwraca się uwagę, jest cena. Dotyczy to oczywiście wielu elementów, ale chyba najbardziej znamienne wydaje się w przypadku złączy. Bezpieczeństwo czy ludzkie życie nie ma swojej ceny, a wymagania stawiane konektorom do zastosowań np. w systemach wojskowych czy kosmicznych są bardzo wysokie. Przyjrzyjmy się dokładniej tym wymogom i spełniającym je elementom.

Złącza to dyskretni, ale absolutnie niezbędni pośrednicy w elektronice. Dzięki nim wszystkie sygnały czy napięcia zasilające trafiają tam, gdzie trzeba. Oczywiście można by zapytać – po co, skoro równie dobrze dałoby się włutować wszystko na stałe... tylko, że nie wszystko w naszym systemie faktycznie powinno być tak połączone. Złącza pozwalają na łatwe dodawanie i odłączanie poszczególnych modułów naszego urządzenia – czy to na etapie produkcji, eksploatacji czy też serwisu urządzenia.

Istnieje wiele sektorów, w których złącza są nad wyraz istotnymi elementami. Wynika to głównie z wymagań w zakresie bezpieczeństwa czy też odporności na warunki środowiskowe. Na ogół to właśnie złącze wystawione jest na zewnątrz urządzenia, poza bezpieczne wnętrze obudowy, więc musi być w stanie poprawnie funkcjonować wszędzie tam, gdzie przewidziano działanie sprzętu. W wielu przypadkach są to naprawdę ekstremalne środowiska. Dodatkowo

często w tych aplikacjach zdarza się, że złącza pełnią rolę nie tylko prostych połączeń serwisowych, ale wręcz strategicznych elementów wpływających na bezpieczeństwo systemu. Awaria takiego komponentu może oznaczać zagrożenie dla bezpieczeństwa czy życia ludzi lub ogromne straty finansowe. W poniższym artykule przyjrzymy się szeregowi sektorów, w których dobór złączy jest bardzo nieprzypadkowy – przeanalizujemy jakie wymagania stawiane są konektorom i zobaczymy przykłady złączy, które te wymagania spełniają... bez pytania o cenę.

Normy ogólne

Konieczność ścisłego opisywania gwarantowanych parametrów złączy (nie tylko tych do zastosowań specjalnych) sprawiła, że powstało wiele norm specyfikujących te parametry. Dla ułatwienia standardy te zebrano w **tabeli 1**. Część z nich dotyczy tylko złączy, ale część norm – IEC 60529 i IEC 60068 – to ogólne specyfikacje odporności urządzeń elektrotechnicznych na warunki środowiskowe, np. wodę oraz pył.

W szczególności warta przypomnienia jest norma **IEC 60529**, specyfikująca tzw. stopień ochrony IP (od ang. ingress protection) urządzenia elektrycznego. Stopień IP mówi o poziomie ochrony urządzenia (przed penetracją czynników zewnętrznych) oraz użytkownika (przed dostępem do niebezpiecznych części urządzenia). W **tabeli 2** zawarto opis poszczególnych klas IP – każda składa się z dwóch cyfr. Pierwsza z nich mówi o ochronie przed dostępem

Tabela 1. Typowe normy stosowane do opisu złączy elektrycznych

Norma	Zakres regulacji	Przykłady zastosowania	Typowe pomiary/badania
IEC 61984	Wymagania bezpieczeństwa i działania złączy elektrycznych i elektronicznych	Wszystkie typy złączy stosowanych w urządzeniach elektrycznych i elektronicznych	Testy bezpieczeństwa elektrycznego, izolacji, ochrony przed dotykiem części czynnych
IEC 60512 (seria)	Metody badań parametrów mechanicznych, elektrycznych i środowiskowych	Testy wytrzymałości, rezystancji styków, trwałości cykli łączenia/rozłączania	Rezystancja styków, siła łączenia/rozłączania, trwałość cykli, testy wibracyjne
IEC 60529	Stopnie ochrony IP przed ciałami stałymi i wodą	Złącza przemysłowe, outdoor, automotive	Testy pyłoszczelności, odporności na zanurzenie, strumień wody
IEC 60068 (seria)	Badania odporności środowiskowej	Złącza dla przemysłu, wojska, lotnictwa	Wibracje, wstrząsy, szok termiczny, wilgotność, mgła solna
IEC 60603-7	Wymiary, parametry i badania złączy modularnych (RJ45, RJ11)	Telekomunikacja, sieci komputerowe	Przesłuchy, tłumienie, rezystancja pętli, integralność sygnału
IEC 60603-2	Wymiary i wymagania dla złączy DIN	Sprzęt audio, przemysł, automatyka	Rezystancja styków, izolacja, dopasowanie mechaniczne

do wnętrza urządzenia przez użytkownika i zanieczyszczenia/pył, a druga opisuje odporność urządzenia na wodę.

Pozostałe normy, które wspomniane będą w dalszej części artykułu, to standardy branżowe, opisujące wymagania stawiane złączom do zastosowań w danej aplikacji.

Złącza dla branży wojskowej (military-grade)

Chyba pierwszym sektorem, jaki przychodzi na myśl, jeśli chodzi o nietypowe, wymagające złącza, są systemy wojskowe. Nie bez powodu – w tej branży wymaga się bezwzględnej niezawodności w niemalże dowolnych warunkach środowiskowych (i to w szerokim zakresie temperatur – nierzadko od -65°C do $+200^{\circ}\text{C}$, włączając w to odporność na ciągłe wibracje, wstrząsy, mgłę solną, promieniowanie UV etc. Wymaga to stosowania specjalnych materiałów:

- **Stopy aluminium** (często z powłoką kadmową lub niklową) – bardzo lekki materiał (gęstość ok. $2,7\text{ g/cm}^3$), ma ogromne znaczenie w konstrukcjach mobilnych, gdzie liczy się dosłownie każdy gram. Jednocześnie aluminium ma dobrą przewodność

cieplną (ułatwia odprowadzanie ciepła ze złącza). Powłoka kadmowa lub niklowa zwiększa odporność na korozję i dodatkowo poprawia przewodność styków.

- **Stal nierdzewna** (np. AISI 316L, 304) okazuje się niezbędna wtedy, gdy wymogiem jest odporność mechaniczna i chemiczna. Ma wysoką wytrzymałość na rozciąganie, jest odporna na działanie mgły solnej i wielu środków chemicznych. Wadą jest waga – około dwa razy większa niż w przypadku aluminium – dlatego stal trafia głównie do złączy stacjonarnych lub narażonych na szczególnie agresywne środowisko.
- **Brąz fosforowy** (fosfobraz) – znany z bardzo dobrych właściwości mechanicznych, szczególnie sprężystości. Cechuje się też bardzo dobrą odpornością na zmęczenie, nie łamie się przy wielu cyklach zginania (nie występuje w nim zjawisko tzw. umocnienia plastycznego, czyli zwiększania twardości – i kruchości – stopu na skutek odkształcenia plastycznego). Dodatkowo stop ten jest bardzo odporny na korozję elektrochemiczną czy sól. W elektronice stanowi jeden z podstawowych materiałów do produkcji styków sprężystych.

Tabela 2. Stopnie ochrony klasy IP (zgodnie z PN-EN 60529:2003)

Cyfra charakt.	Pierwsza cyfra (stopień ochrony przed pyłem)	Dru ga cyfra (stopień ochrony przed wodą)
0	bez ochrony	bez ochrony
1	ochrona przed dostępem do części niebezpiecznych wierzchem dłoni ochrona przed obcymi ciałami stałymi o średnicy $\geq 50\text{ mm}$	ochrona przed padającymi kroplami wody
2	ochrona przed dostępem do części niebezpiecznych palcem ochrona przed obcymi ciałami stałymi o średnicy $\geq 12,5\text{ mm}$	ochrona przed padającymi kroplami wody przy wychyleniu obudowy o dowolny kąt do 15° od pionu w każdą stronę
3	ochrona przed dostępem do części niebezpiecznych narzędziem ochrona przed obcymi ciałami stałymi o średnicy $\geq 2,5\text{ mm}$	ochrona przed natryskiwaniem wodą pod dowolnym kątem do 60° od pionu z każdej strony
4	ochrona przed dostępem do części niebezpiecznych drutem ochrona przed obcymi ciałami stałymi o średnicy $\geq 1\text{ mm}$	ochrona przed bryzgami wody z dowolnego kierunku
5	ochrona przed dostępem do części niebezpiecznych drutem ochrona przed pyłem	ochrona przed strugą wody ($12,5\text{ l/min}$) laną na obudowę z dowolnej strony
6	ochrona przed dostępem do części niebezpiecznych drutem ochrona pyłoszczelna	ochrona przed silną strugą wody (100 l/min) laną na obudowę z dowolnej strony
7		ochrona przed skutkami krótkiego zanurzenia w wodzie (30 min na głębokość $0,15\text{ m}$ ponad wierzchem obudowy lub 1 m od spodu dla obudów niższych niż $0,85\text{ m}$)
8		ochrona przed skutkami ciągłego zanurzenia w wodzie (obudowa ciągle zanurzona w wodzie, w warunkach uzgodnionych między producentem i użytkownikiem, lecz surowszych niż według cyfry 7)
9		ochrona przed zalaniem silną strugą wody pod ciśnieniem ($8\ldots 10\text{ MPa}$ i temp. 80°C) zgodnie z normą DIN 40050 (oznaczenie to nie występuje w PN-EN 60529:2003)

- **Tytan i jego stopy** to materiały o ekstremalnej odporności na korozję (np. w środowiskach z obecnością paliw lub materiałów żrących). Stopy te łączą w sobie wysoką odporność chemiczną z lekkością (gęstość tylko nieco większa niż aluminium, ale znacznie niższa niż w przypadku stali). Materiał ten jest całkowicie niemagnetyczny.
- **Mosiądz niklowany/cynowany** – materiał o bardzo dobrej przewodności elektrycznej i względnie niskiej cenie. Używa się go jako bazowego materiału na styki i kontakty w złączach. Powłoki (niklowa, cynowa lub srebrna) redukuje rezystancję styku oraz zabezpiecza go przed utlenianiem.
- **Miedź berylowa (CuBe)** – to materiał na styki o bardzo dużej sprężystości i odporności na przepalenie, stosowany w połączeniach narażonych na wysokie prądy uderowe lub wielokrotne łączenie. Miedź berylowa zachowuje parametry w szerokim zakresie temperatur (do ok. 200...250°C), co odpowiada wymaganiom militarnym.
- **PEEK (Polyether Ether Ketone)** to termoplastyczny polimer o wyjątkowej odporności na temperaturę (do 260°C) i promieniowanie (UV i jonizujące). Ma on niski współczynnik absorpcji wilgoci, nie emituje gazów (co jest ważne w hermetycznych urządzeniach) i charakteryzuje się wyjątkową odpornością chemiczną. Stosowany jest w systemach elektronicznych w rakietach czy radiostacjach polowych.
- **PTFE (Teflon)** to polimer chętnie stosowany jako izolator w złączach RF. Cechuje się ekstremalnie niskim współczynnikiem strat dielektrycznych oraz niemalże całkowitą odpornością chemiczną. Praktycznie nie starzeje się i zachowuje parametry w szerokim zakresie temperatur (od -200°C do 260°C).
- **Elastomery fluorowe** (np. FKM/Viton, EPDM) są głównie wykorzystywane do produkcji uszczelek i O-ringów. W zależności od typu mogą być odporne na paliwa, oleje, wilgoć, płyny hydrauliczne czy ozon. Elementy wykonane z tych materiałów są kluczowe dla zapewnienia wysokiej klasy IP w wymagających warunkach środowiskowych.
- **Kompozyty polimerowe** (np. PPS, LCP z włóknem szklanym) są często używane tam, gdzie konieczne jest obniżenie masy bez utraty sztywności mechanicznej. Kompozyty mogą mieć podobną wytrzymałość do aluminium przy masie niższej o 30...40%. Często stosowane są w złączach do płytek drukowanych.
- **Żywicę epoksydowe** są najczęściej stosowane jako zalewy ze względu na bardzo wysoką wytrzymałość mechaniczną i stabilność wymiarową – po utwardzeniu tworzą monolity, które nie zmieniają swoich właściwości w temperaturach od -55°C do 180°C. Charakteryzują się również wysoką odpornością na oleje, paliwa i inne środki.
- **Żywicę poliuretanowe** wybiera się tam, gdzie oprócz ochrony wymagane jest również tłumienie drgań – są bardziej elastyczne i wibrują razem z przewodem. Są też nieco bardziej odporne na szok termiczny, dlatego stosuje się je np. w miejscach narażonych na bardzo gwałtowne zmiany temperatury.
- **Żywicę silikonowe** znajdują zastosowanie rzadziej, głównie w złączach o bardzo rozbudowanej części tylnej, z dużą ilością przewodów – ponieważ mają bardzo niską lepkość i łatwo wypełniają nawet przedłużone kanały. Ich największą zaletą jest odporność na promieniowanie UV oraz zachowanie elastyczności nawet poniżej -60°C, dlatego czasem stosuje się je w systemach rozmieszczonych na zewnętrznych powierzchniach pojazdów lub w antenach.

W systemach wojskowych standaryzacja i interoperacyjność (zdolność komponentów – pochodzących nawet od różnych producentów – do współdziałania ze sobą w sposób bezproblemowy, bez konieczności dodatkowych modyfikacji) są kluczowe. Z tego powodu wiele złączy zdefiniowanych jest za pomocą norm: MIL-DTL-38999, MIL-DTL-26482, MIL-DTL-5015, VG95234 i innych. Norma taka

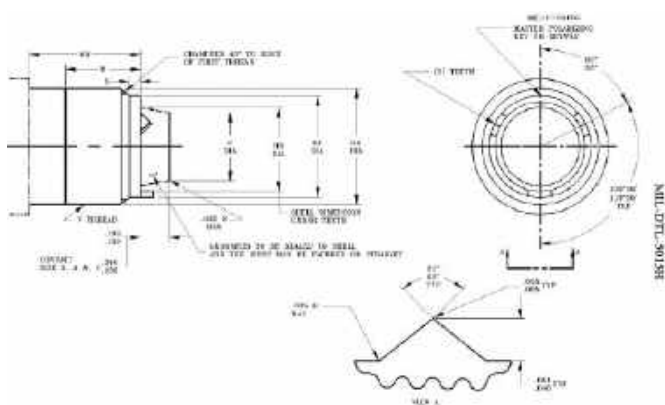


Fotografia 1. Złącze typu MIL-DTL-26482 produkcji Aero Electric (<https://www.prowireusa.com/>) ma aluminiową, niklowaną obudowę i złocionymi pinami z fosfobrazu. Dostępna jest również wersja w obudowie z pokryciem kadmowym lub anodyzowana na czarno

specyfikuje nie tylko ogólne parametry mechaniczne i elektryczne, ale również projekt mechaniczny, rozkład i wielkość pinów oraz inne detale konstrukcji złącza. Na **rysunku 1** pokazano przykładowy fragment takiej normy (MIL-DTL-5015), który opisuje konstrukcję dławnicy kablowej na przewodzie, mającej zapewnić poprawne jego utrzymanie we wtyczce. Oprócz detali konstrukcji mechanicznej w normie zapisano też wiele innych parametrów – odporność mechaniczną (na uderzenia czy na wibracje), sposób galwanicznego pokrycia



Fotografia 2. Złącza z rodziny MIL-DTL-38999 firmy Souriau (<https://meds-int.com/>)



Rysunek 1. Norma MIL-DTL-5015H z 18 maja 2000 r. (fragment)

styków, materiał i sposób zalania gotowego złącza – powołując się przy tym na inne normy wojskowe.

Jednym z zabiegów często stosowanych w złączach wojskowych (i nie tylko) jest zalewanie wnętrza złącza – miejsca wyprowadzenia styków oraz połączenia przewodów. Obszar ten zostaje całkowicie wypełniony specjalną masą zalewową. Takie rozwiązanie stosuje się przede wszystkim w złączach narażonych na skrajne warunki środowiskowe, w tym ciągłe działanie wilgoci, błota, mgły solnej czy wibracji, gdzie standardowa uszczelka może z czasem przestać pełnić swoją funkcję. Zalewaniu poddaje się zarówno złącza okrągłe (np. w standardzie MIL-DTL-38999) jak i prostokątne, montowane na PCB lub na przewodach – zwłaszcza wtedy, gdy złącze pozostaje na stałe w systemie i nie ma potrzeby okresowego wymieniania go (zwłaszcza w warunkach polowych).

Oprócz wspomnianych wcześniej żywic w niektórych konstrukcjach stosuje się materiały termoplastyczne nakładane na gorąco. Taka technologia jest wykorzystywana do zalewania tylnej części złącza przy pomocy termoplastu, podawanego pod ciśnieniem w stanie uplastycznionym lub jako tzw. „overmoulding” całego złącza. Ta druga technologia polega na utworzeniu dodatkowej, zewnętrznej powłoki, formowanej wtryskowo na gotowym złączu i jego przewodach (co pozwala na utworzenie odgiętki), jak pokazano na **fotografii 4**. Stosuje się tutaj najczęściej poliamidy (PA6, PA12) z dodatkiem włókna szklanego, PPS albo elastyczne termoplasty poliuretanowe (TPU). Nakładane są one w formie dwuetapowej: najpierw montuje się złącze i przewody, a następnie całość trafia do formy, w której tworzy się monolityczną powłokę.

Podczas procesu zalewania niezwykle istotne jest całkowite odgazowanie zalewy żywicznej oraz dokładne ułożenie przewodów w prowadnicach, aby nie powstały puste przestrzenie lub nadmierne naprężenia na pinach czy oczkach lutowniczych złącza. Po utwardzeniu masa pełni jednocześnie funkcję izolatora elektrycznego, uszczelnienia i mechanicznego mocowania przewodów, dzięki czemu złącze zachowuje pełną funkcjonalność nawet po wielogodzinnej ekspozycji na wstrząsy, nacisk czy niekorzystne warunki środowiskowe.

W przypadku termoplastycznego overmouldingu dodatkowo zapewnia on, oprócz izolacji elektrycznej i ochrony złącza przed czynnikami środowiskowymi, wyjątkową odporność na zginanie przewodu tuż przy wyjściu ze złącza (praktycznie eliminując problem urwania żyły przy samym korpusie). Projektowanie takiej dodatkowej ochrony złącza to skomplikowany proces i na ogół stosuje się standardowe overmouldingi dla danego rodzaju komponentu. Niestandardowe rozwiązania mogą mieć np. dodatkowe otwory montażowe, integrować w sobie kilka złączy etc. – tak, aby dostosować element do konkretnej aplikacji. Należy jednak pamiętać, że zlecenie wykonywania niestandardowego overmouldingu oznacza konieczność wyprodukowania formy wtryskowej, co wiąże się na ogół z bardzo wysokimi kosztami i ma uzasadnienie jedynie w przypadku dużych serii (lub dużych budżetów).



Fotografia 3. Okrągłe złącze wojskowe typu MIL-DTL-5015 zalane żywicą w tylnej części. (Źródło: strona użytkownika AlexDG1993 na portalu Reddit)



Fotografia 4. Overmoulding – zalewanie złącza i przewodu

Ostatnim aspektem, o jakim należy powiedzieć w kontekście złączy stosowanych w systemach wojskowych, są zagadnienia związane z tzw. inżynierią użyteczności. Przy projektowaniu złączy do zastosowań wojskowych przykładą się do tego bardzo dużą wagę. Operator musi być w stanie szybko i bezbłędnie podłączyć odpowiedni moduł, nawet przy ograniczonej widoczności, w rękawicach lub pod presją czasu. Z tego powodu złącza są wyposażone w różne systemy kodowania mechanicznego (klucze), które uniemożliwiają wpięcie złącza w nieodpowiednie miejsce. Oprócz klasycznych kluczy rowkowych stosuje się również kodowanie poprzez różny kąt ustawienia pinów prowadzących albo asymetryczne ułożenie śrub mocujących. W praktyce oznacza to, że nawet jeżeli kilka złączy wygląda z zewnątrz identycznie, każde można podpiąć tylko do właściwego gniazda. Jednocześnie producenci stosują wyrażne oznaczenia w formie kolorowych wkładek, pierścieni lub znaków grawerowanych, które są spójne z oznaczeniami stosowanymi typowo w wojskach NATO (np. granatowy dla zasilania, żółty dla sygnałów optoelektronicznych, czerwony dla linii sterowania). Należy o tych aspektach pamiętać, dobierając odpowiednie złącza do projektowanych przez nas systemów.

Opisane powyżej złącza i normy korzystają z amerykańskich standardów MIL-STD – nie są to jedyne tego rodzaju normy na świecie. W Polsce istnieją normy obronne (NO), Rosjanie korzystają z norm GOST (GOST); ich warianty wojskowe oznaczone są jako OCT (OST) lub GOST B (GOST V). Te pierwsze specyfikują bardziej szczegółowe standardy konstrukcji i testów środowiskowych; normy OCT B oraz OCT 4 to niemalże odpowiedniki MIL-STD-1344 i IEC 60068. Normy wojskowe GOST B, obejmujące złącza okrągłe, prostokątne, wielopinowe, często odpowiadają funkcjonalnie serii MIL-DTL-38999, ale oferują pełnej kompatybilności z nimi (złącza te mają inny układ kluczy i rozstaw pinów). Dodatkowo część rosyjskich konstrukcji bazuje na starych normach PM (ReMo) – typowych dla produkcji lotniczej z czasów ZSRR. W Chinach wykorzystywane są standardy opublikowane przez urząd normalizacyjny GB (Guobiao) oraz normy wojskowe GJB (GuojunBiao). Jeśli chodzi o złącza do zastosowań militarnych podstawowymi są GJB 598, opisująca ogólną specyfikację złączy okrągłych dla przemysłu zbrojeniowego (nieformalnie jest to chiński odpowiednik złączy MIL-DTL-38999), GJB



Fotografia 5. Złącze koncentryczne BNC firmy Gigaflight wraz z akcesoriami i koncentrycznym przewodem 50 Ω w wykonaniu lotniczym (<https://www.gigaflightinc.com/>)

Tabela 3. Normy STANAG dotyczące złączy

STANAG	Zakres normy	Opis	Typowe kompatybilne złącza MIL
STANAG 1008	Zasilanie 28 V DC w pojazdach wojskowych	Definiuje mechaniczny układ wtyków i gniazd zasilających oraz układ pinów we wtyku	MIL-DTL-26482 (seria I/II), MIL-DTL-5015 (okrągłe złącza zasilające)
STANAG 4007	Zasilanie 115/200 V AC, 400 Hz w statkach powietrznych	Określa konstrukcję złączy zasilania pokładowego do zastosowań w statkach powietrznych	MIL-DTL-38999 (Seria II i III w wersji wysokoprądowej)
STANAG 4074	Interfejsy elektryczne dla uzbrojenia podwieszanego	Opisuje rodzaj i wymagania co do wielopinowych złączy o wysokiej odporności na wibracje, stosowanych do uzbrojenia podwieszanego na statkach powietrznych	MIL-DTL-38999 (Seria III), opcjonalnie MIL-DTL-83723
STANAG 3681	Przenośne sieci łączności taktycznej	Definiuje wymagania dla złączy, które mają być odporne na wodę, pył i błoto	MIL-DTL-55116 (złącza audio), MIL-DTL-38999 (złącza miniaturowe)
STANAG 3734	Systemy transmisji danych w pojazdach lądowych	Opisuje wymagania dla złączy oraz interfejsy (elektryczne i logiczne) dla magistral stosowanych w pojazdach lądowych (CAN czy MIL-STD-1553)	MIL-DTL-38999 (z wkładką do transmisji danych), MIL-DTL-83513 (złącza typu Micro-D)

101A (opisująca metody badań środowiskowych i mechanicznych – odpowiednik norm MIL-STD-202 oraz IEC 60068), GJB 376 (opisuje miniaturowe złącza do sprzętu lotniczego i rakietowego) oraz GJB 1428 (standard interfejsów do mobilnych urządzeń łączności, takich jak np. terminale polowe). W praktyce wiele złączy wojskowych stosowanych w Chinach jest kompatybilnych ze złączami MIL-DTL-38999 czy MIL-DTL-26482, ale produkowane są one zgodnie z normami GJB i mają własny zakres procedur kwalifikacyjnych (np. dodatkowe badania EMC lub testy pracy w atmosferze z pyłem kwarcowym).

Normy STANAG (normy obejmujące cały sojusz NATO) odnoszące się do złączy nie tworzą jednego spójnego katalogu, ale pojawiają się jako część konkretnych standaryzacji dotyczących zasilania, transmisji danych albo interfejsów uzbrojenia. W tabeli 3 opisano najczęściej spotykane normy STANAG dotyczące złączy oraz typowe złącza wojskowe, które spełniają ich wymagania.

Na rynku dostępnych jest wiele złączy spełniających wojskowe normy. Z uwagi na fakt, że wiele z nich zdefiniowane jest w tych normach, złącza różnych producentów są ze sobą kompatybilne. Typowymi firmami, które produkują komponenty dla opisywanego segmentu, są firmy takie jak: Amphenol, TE Connectivity, Souriau-Sunbank, ITT Cannon czy Glenair.

Złącza wojskowe łączą w sobie chyba wszystkie aspekty związane z pracą w trudnych środowiskach, a jednocześnie wiele z nich stosowanych jest relatywnie masowo. Materiały i technologie używane w tym sektorze będą przewijały się również w innych branżach, wymagających złączy odpornych na trudne warunki środowiskowe itp. Co więcej, wiele sektorów rynku korzysta po prostu ze złączy

wojskowych, bo są one już przetestowane w realnych warunkach i tworzenie nowych specyfikacji nie miałoby sensu.

Złącza dla lotnictwa

W przypadku złączy lotniczych kluczowy aspekt to odporność na silne wibracje i drgania występujące podczas lotu. Połączenie nie może się poluzować ani zmieniać impedancji, nawet przy długotrwałej pracy silnika lub podczas manewrów z dużym przeciążeniem. Drugim istotnym zagadnieniem jest masa – każdy dodatkowy gram przekłada się na zwiększone zużycie paliwa, stąd stosowanie stopów aluminium, tytanu oraz zminiaturyzowanych układów pinów są tak często spotykane w tym sektorze.

Osobnym aspektem jest stabilność pracy w szerokim zakresie temperatur, typowo od -55°C do 175°C , przy czym zmiany temperatury mogą następować bardzo gwałtownie (zwłaszcza w czasie startu). W pewnych obszarach wymagania te mogą być znacznie wyższe, na przykład złącza pracujące w obrębie silników odrzutowych podlegają jeszcze bardziej rygorystycznym wymaganiom, w zależności od miejsca montażu w otoczeniu napędu. W sekcji sprężarki złącza muszą pracować w temperaturach ciągłych rzędu $175\ldots 200^{\circ}\text{C}$, a przy samej sekcji gorącej (komora spalania oraz turbina) przyjmuje się zakres temperatur dochodzący do 260°C , a w niektórych aplikacjach nawet do 300°C . To, co kluczowe w tym aspekcie, to nie tylko możliwość chwilowej pracy w tej temperaturze, ale także odporność na długotrwałe przeciążenia termiczne, w połączeniu z silnymi wibracjami i cyklami grzania/chłodzenia. Na przykład złącza firmy Apollo z serii HTX (fotografia 6) zaprojektowane są do pracy w temperaturach do 440°C . Oprócz tego są one w pełni



Fotografia 6. Złącza z serii HTX firmy Apollo, zaprojektowane do pracy w temperaturach do 440°C w silnikach odrzutowych (<https://www.apollo-aerospace.com>)



Fotografia 7. Wiązki kablowe do samolotów (<https://dougselectrical.com>)



Fotografia 8. Izolująca, elastyczna taśma, służąca do owijania fragmentów złączy i połączeń we wiązkach kablowych (<https://www.avdec.com>)

hermetyczne i odporne na szeroki zakres czynników środowiskowych oraz wibracje.

Istotną okazuje również odporność na warunki atmosferyczne: kondensację wilgoci, zmianę ciśnienia, mgłę solną oraz promieniowanie UV. Dodatkowo złącza w systemach krytycznych (awionika, systemy sterowania lotem, czujniki silnika) muszą zapewniać bardzo niską rezystancję kontaktową, wysoką odporność na zakłócenia EMI/RFI i implementować mechaniczne zabezpieczenie przed przypadkowym rozłączeniem – np. bagietkowe z zabezpieczeniem, klips 2-stopniowy lub śrubę o znanym momencie dokręcania.

Oprócz odporności na temperaturę, wibrację etc. wymagany jest również bardzo niski współczynnik odgazowywania materiałów, z których wykonano złącze, ponieważ w środowisku o wysokiej temperaturze i często niskim ciśnieniu, jakkolwiek migracja gazów lub związków organicznych mogłaby spowodować zanieczyszczenie sensorów lub zaburzenie funkcjonalności systemu. Dlatego też w złączach tych stosuje się wkładki izolacyjne wykonane z PEEK, ceramiki lub spieków szklanych (jak w złączach Apollo HTX). Obudowy tych elementów wykonywane są często z tytanu albo niklowanych stopów stali, a styki pokrywa się złotem (co zapewnia stabilny kontakt nawet przy wysokiej temperaturze).

Duże znaczenie w systemach awionicznych mają także wiązki kablowe. Podobnie jak w aplikacjach motoryzacyjnych, w samolotach znajdują się kilometry okablowania, które przygotowywane jest często niezależnie od budowy samego samolotu – a następnie instalowane na pokładzie. Przykład takiej wiązki pokazano na **fotografii 7**.

Integracja przewodów i złączy we wiązki stawia dodatkowe wymagania konektorom. Muszą one pozwalać nie tylko na odpowiednie zamocowanie przewodów, ale często również dodatkowych peszli ochronnych czy oplotów do przewodów. Peszle te często są dodatkowo izolowane, zapewniając zapasową barierę przed wnikiem wilgoci do wnętrza wiązek kablowych i złączy. Stosuje się często dodatkowe uszczelnienia, takie jak kity czy pokazana na **fotografii 8** taśma uszczelniająca, zapewniająca dodatkową ochronę przed wnikiem wilgoci i korozją.

Jeśli chodzi o normy specyfikujące złącza stosowane w sektorze lotniczym, to – oprócz ogólnych norm i części norm wojskowych opisanych powyżej – zastosowania mają następujące normy:

- EN 2997 – europejska norma dot. okrągłych złączy lotniczych, stosowanych w układach elektrycznych i sterujących,
- EN 3155 – norma opisująca styki elektryczne używane w złączach lotniczych (w tym styki zaciskane i tzw. high-density),
- ARINC 600 i ARINC 404 – standardy interfejsowe dla modułowych złączy stosowanych w awionice pokładowej,

- RTCA DO-160 – norma cywilna używana przez FAA (amerykańska Federalna Administracja Lotnictwa) oraz EASA (Agencja Unii Europejskiej ds. Bezpieczeństwa Lotniczego), określająca odporność elementów elektronicznych (w tym złączy) na temperaturę, ciśnienie, promieniowanie, wibracje i zakłócenia elektromagnetyczne w systemach lotniczych.

Na rynku lotniczym działa stosunkowo wąska, ale bardzo wyspecjalizowana grupa producentów złączy. Amphenol Aerospace oraz TE Connectivity (DEUTSCH/Raychem) należą do największych i oferują szeroką gamę złączy, głównie stosowanych w awionice, sterowaniu lotem czy czujnikach silnikowych (w tym serie kwalifikowane do pracy do 260°C). Mniejsi producenci, jak np. Apollo Aerospace, produkują często złącza o bardziej niszowych zastosowaniach, jak np. elementy do pracy w samych silnikach, przewidziane do temperatur do 440°C. Firma Souriau-Sunbank (obecnie część Eaton) jest jednym z głównych dostawców złączy okrągłych. Radiall specjalizuje się w złączach RF/mikrofalowych do instalacji komunikacyjnych i sensorów radarowych, natomiast Smiths Interconnect dostarcza zaawansowane złącza wysokotemperaturowe i hermetyczne do gorących sekcji silników odrzutowych. Glenair produkuje m.in. miniaturowe złącza do systemów sterowania i czujników paliwowych, a z kolei firmy LEMO oraz Fischer Connectors są dobrze znane z serii lekkich i szczelnych złączy przeznaczonych do zastosowań awionicznych o ograniczonej przestrzeni montażowej. Część tych marek wymieniana była również przy okazji omawiania złączy wojskowych.

Złącza do systemów kosmicznych

Złącza dla kosmonautyki mają podobne wymagania, jak komponenty lotnicze, z kilkoma dodatkowymi wymogami, wynikającymi z pracy poza ziemską atmosferą. Przede wszystkim muszą bezawaryjnie działać w próżni co oznacza, że wszystkie użyte materiały – zarówno metale, jak i tworzywa sztuczne – muszą mieć bardzo niski współczynnik odgazowania (ang. low outgassing). Odgazowanie to uwalnianie gazu, który został rozpuszczony, uwięziony, zamrożony lub zaabsorbowany w jakimś materiale. Proces ten może obejmować sublimację jak, i parowanie (które są przejściami fazowymi substancji w gaz), a także desorpcję, przesiąkanie z pęknięć lub objętości wewnętrznych oraz gazowe produkty powolnych reakcji chemicznych, zachodzących w materiale. Wrzenie jest zazwyczaj uważane za zjawisko odrębne od odgazowania, ponieważ polega na przejściu fazowym cieczy w parę tej samej substancji. Niezależnie od tego, każde wydzielanie lotnych substancji w warunkach próżni może prowadzić do ich kondensacji na powierzchni, co okazuje się problematyczne zwłaszcza w przypadku czułych elementów w systemie kosmicznym: optyki, sensorów lub paneli słonecznych. Gromadzenie nalotów na powierzchniach krytycznych może z kolei spowodować trwałe uszkodzenie systemu i stanowi zagrożenie dla powodzenia całej misji. Z tego powodu tylko wąska



Fotografia 9. Złącza wielopinowe przeznaczone do pracy w wysokiej próżni, wykonane z materiałów o niskim współczynniku odgazowania (<https://globetech.jp/>)



Fotografia 10. Hermetyczne złącza Micro-D, zapewniające próżniowo-szczelne połączenia przez ściany urządzenia (<https://douglaselectrical.com/>)

grupa materiałów, np. PEEK, PTFE, ceramika czy stopy niklu lub tytanu, dopuszczone są do zastosowań w przestrzeni kosmicznej.

Co ciekawe, podobne wymagania co do odgazowywania stawiane są jeszcze w innej branży, chyba jeszcze mniej znanej szerokiego gronu inżynierów – mowa o systemach wysokiej próżni, które można spotkać np. w fabrykach półprzewodników, napylarkach czy instalacjach akceleratorów cząstek czy komputerach kwantowych. W tej grupie aplikacji element nie może emitować żadnych gazów, ponieważ mogłoby to zanieczyścić próżnię systemu i zwiększyć panujące w niej ciśnienie.

Kolejnym kluczowym czynnikiem jest odporność na wysoki poziom promieniowania jonizującego – zarówno promieniowanie kosmiczne, jak i wysokoenergetyczne cząstki pochodzące z wiatru słonecznego, mogą powodować stopniową degradację materiałów organicznych oraz zmianę ich parametrów elektrycznych. W praktyce oznacza to konieczność stosowania powłok metalicznych na stykach (stosuje się złoto, pallad, srebro) oraz wyeliminowania materiałów szczególnie podatnych na uszkodzenia radiacyjne (np. klasycznych epoksydów).

Złącza w aplikacjach kosmicznych pracują też w warunkach bardzo dużych wahań temperatury: często od -150°C (od strony zaciemnionej) do $+200^{\circ}\text{C}$ (od strony nasłonecznionej), przy czym zmiany pomiędzy obydwojoma tymi położeniami mogą następować w bardzo krótkim czasie. Konstrukcja złącza musi gwarantować, że nawet po kilkuset takich cyklach nie dojdzie do np. mikropęknięć obudowy czy jej rozszczelnienia. Dotyczy to także pinów złącza, które nie mogą tracić sprężystości w takich warunkach, gdyż mogłoby się to przełożyć się na pogorszenie styku, zwiększenie rezystancji połączenia lub zmianę impedancji charakterystycznej.

Bardzo istotnym czynnikiem w przypadku złączy do aplikacji kosmicznych jest również odporność na drgania i uderzenia mechaniczne występujące podczas startu rakiety. Zanim złącze trafi w stabilne środowisko orbitalne, musi wytrzymać znacznie większe obciążenia niż te, które pojawiają się w eksploatacji na orbicie. Obciążenia/wibracje są istotnie większe, niż występujące w systemach lotniczych, ale trwają znacznie krócej. Z tego względu stosuje się mechaniczne blokady do złączy w postaci zdwojonych gwintów, pierścieni blokujących czy zamków bagnetowych z zabezpieczeniem.

Podobnie jak w przypadku systemów lotniczych, w branży kosmicznej nie wolno pominąć kwestii miniaturyzacji i masy – każdy gram wynoszonego wyposażenia to realny koszt (i to niemały), dlatego złącza kosmiczne są ekstremalnie lekkie, a konfiguracje styków – zoptymalizowane tak, by zajmowały minimalną przestrzeń (stąd popularność złączy z serii Micro-D i Nano-D).

Tak jak w przypadku systemów lotniczych, tak i w pojazdach kosmicznych czy satelitach wiązki kablowe są niezwykle

często i chętnie stosowane przez konstruktorów. Z uwagi jednak na ograniczenia materiałowe, istnieje o wiele mniejszy wybór materiałów do ich konstrukcji. Stosuje się często metalowe peszle, które chronią przewody przed uszkodzeniem, np. przetarciem na skutek wibracji, a także ekranują je przed zakłóceniami elektromagnetycznymi (**fotografia 11**). Wszystkie elementy takiej wiązki muszą być wykonane z materiałów kompatybilnych z próżnią, stąd też nawet spinki, utrzymujące peszle na swoim miejscu, muszą być metalowe.

W przypadku systemów kosmicznych stosuje się zestaw norm dedykowanych dla przemysłu space, które są odrębne od klasycznych standardów lotniczych czy wojskowych. Najważniejsze z nich to:

- Normy ECSS (European Cooperation for Space Standardization), stosowane w europejskich programach ESA:
 - ECSS-Q-ST-30 – ogólne wymagania jakościowe dla komponentów elektrycznych (obejmuje złącza),
 - ECSS-Q-ST-30-11 – szczegółowe wymagania dot. złączy i ich akcesoriów, stosowanych w systemach kosmicznych. Standard definiuje m.in. odporność na radiację, wymagany poziom odgazowywania i metody badań w próżni,
 - ECSS-Q-ST-70 – opis techniki lutowania i łączenia przewodów (znacząca część poświęcona jest terminacji styków w złączach),
 - ECSS-Q-ST-20-07 – norma opisująca metody kwalifikacji i testów środowiskowych komponentów (m.in. testy termiczne, odporności na wibracje, pracę w próżni czy uderzenia mechaniczne).
- NASA-STD-8739.4 (Crimping, Interconnecting Cables, Harnesses, and Wiring) – norma ta obowiązuje w programach NASA, opisuje szczegółowo wymagania dla styków złączy, kabli i sposobu ich łączenia.
- MIL-DTL-83513/MIL-DTL-32139 – normy wojskowe specyfikujące m.in. złącza micro-D, które są powszechnie stosowane w systemach kosmicznych jako bazowe standardy dla mikro- i nanozłączy, dlatego często są wpisywane w specyfikację komponentów kosmicznych jako „standard bazowy”.

Dodatkowo w kontraktach kosmicznych pojawiają się wymagania dotyczące odgazowania materiałów. Specyfikuje się je zgodnie z normami NASA ASTM-E595 lub ECSS-Q-ST-70-02, które bezpośrednio odnoszą się do materiałów używanych w złączach (wkładki dielektryczne, powłoki styków etc).

W segmencie technologii kosmicznych działa stosunkowo mała grupa producentów, ponieważ wymagania związane z kwalifikacją do przestrzeni kosmicznej są bardzo restrykcyjne i kosztowne. Do najważniejszych firm należą Axon Cable (Francja), jeden z głównych dostawców mikro- i nanozłączy dopuszczonych tak przez ESA, jak i NASA (serie Micro-D i Nano-D oraz złącza do paneli słonecznych).



Fotografia 11. Złącza Micro-D firmy Axon, wyposażone w szybkozłącza z dodatkowymi zapięciami, zabezpieczające przed rozłączeniem na skutek wibracji. Widoczne są także metalowe peszle kablowe z (również metalowymi) spinkami, zabezpieczające przewody przed uszkodzeniem oraz EMI (<https://www.axon-cable.com/>)



Fotografia 12. Przykładowe złącza do instalacji podmorskich firmy CRE (<https://www.cre-marine.com/>)

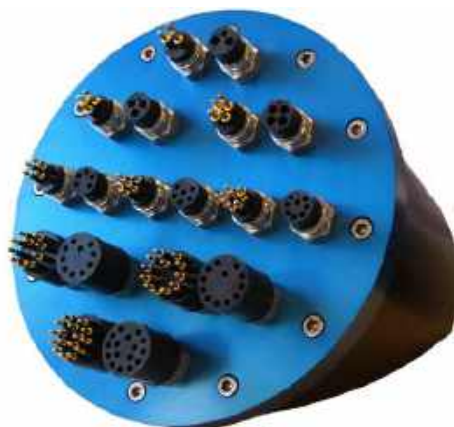
Smiths Interconnect dostarcza hermetyczne złącza wysokotemperaturowe oraz RF, stosowane w satelitach i sondach kosmicznych. Firma Glenair ma w ofercie całą gamę lekkich złączy space-grade (w tym wersje specjalnie przygotowane do pracy w próżni i z ograniczonym odgazowaniem). TE Connectivity (DEUTSCH) produkuje szereg złączy o gęstym rastrze do satelitów i rakiet nośnych, w tym bardzo popularne serie NanoRF do transmisji sygnałów RF. Radiall dostarcza złącza RF (w tym mikrofalowe) oraz komponenty przeznaczone do optycznych łączy światłowodowych (na przykład między modułami satelitów). Warto wymienić również Harwin – produkowane przez firmę serie Gecko (1,25 mm i 2,0 mm) zostały opracowane specjalnie z myślą o wymaganiach sektora „New Space” i znajdują zastosowanie w platformach typu Cubesat, gdzie wymagania co do parametrów niezawodnościowych są zmniejszone względem klasycznych, „dużych” systemów kosmicznych.

Złącza do przemysłu morskiego

Środowisko morskie, wbrew pozorom, jest nie mniej wymagające niż przestrzeń kosmiczna, jednakże ma zupełnie odmienny zestaw wymagań. Złącza przeznaczone do pracy w wodzie muszą być przede wszystkim odporne na korozję elektrochemiczną, ponieważ kontakt z wodą morską i mgłą solną zapewnia wyjątkowo agresywne warunki chemiczne. Ta odporność musi być utrzymana nie tylko podczas krótkotrwałego zanurzenia – w wielu aplikacjach (sonary, ROV, czujniki kadłubowe) złącze przebywa ciągle pod wodą, często na dużej głębokości, gdzie ciśnienie hydrostatyczne sięga dziesiątek barów. Dodatkowo wymaga się pełnej szczelności IP68/IP69K oraz stabilności mechanicznej przy stałym naporze wody. W instalacjach pokładowych złącza narażone są również na intensywne wibracje i zmiany temperatury (od -40°C do $+85...100^{\circ}\text{C}$), dlatego muszą zachowywać parametry styków również pod dynamicznym obciążeniem.

Ze względu na wymaganie odporności na korozję, do produkcji obudów złączy morskich stosuje się najczęściej miedzionikiel (CuNi10Fe1Mn), charakteryzujący się wyjątkową odpornością korozyjną oraz praktycznie nie wykazujący podatności na tworzenie ogniw galwanicznych na styku z innymi, typowymi materiałami okrętowymi. W złączach podwodnych stosuje się również stale duplex i superduplex, które mają dużo wyższą wytrzymałość mechaniczną od zwykłej stali nierdzewnej oraz znacznie większą odporność na korozję i pękanie. Do pracy na większych głębokościach stosuje się także tytan, który łączy niską masę z niemal całkowitą odpornością na wodę morską i brakiem magnetyczności (co jest istotne przy systemach sonarowych).

W elementach uszczelniających złączy stosowane są elastomery fluorowe (np. Viton, FKM) i EPDM, które zachowują elastyczność przy niskich temperaturach i wykazują odporność na działanie soli, olejów oraz paliw okrętowych. Styki wykonuje się zazwyczaj z brązu berylowego lub miedzi z powłoką złotą, co zapewnia bardzo niską rezystancję kontaktową oraz zapobiega utlenianiu i degradacji powierzchni roboczej.



Fotografia 13. Rodzina złączy GLMC firmy Glenair, przeznaczonych do pracy podwodnej na głębokości do 6800 metrów (<https://www.technologycatalogue.com/>)

W przypadku systemów morskich nie istnieje jedna uniwersalna norma dla złączy, ponieważ obowiązują normy okrętowe, definiujące warunki środowiskowe, wymagania materiałowe i metody badań – a producenci złączy muszą wykazać ich spełnienie. Najważniejsze z nich to:

- IEC 60092 – zestaw międzynarodowych norm poświęcony instalacjom elektrycznym na statkach. Zawiera m.in. wymagania dla złączy zasilających, sygnałowych i sterujących, a także klasy ochrony IP i odporność na mgłę solną,
- zasady ogólne DNV-i norma DNV-CG-0339 – wymagania klasyfikacyjne norweskiego towarzystwa DNV dla komponentów elektrycznych, określające odporność na drgania, korozję i ciśnienie hydrostatyczne (złącza muszą przejść odpowiednie kwalifikacje, aby być dopuszczone do użycia na jednostkach pływających),
- rekomendacje Lloyd's Register obejmują testy środowiskowe dla złączy stosowanych w instalacjach pokładowych (odporność na wibracje, uderzenia, mgłę solną i cykliczne zmiany temperatury),
- MIL-STD-810, metoda 509, specyfikująca metodologię pomiaru odporności na mgłę solną. Choć jest to norma wojskowa, została powszechnie zaadaptowana jako punkt odniesienia w morskich aplikacjach cywilnych,
- API RP 17F (Offshore Subsea Electrical Power and Control Systems) – norma stosowana do opisu budowy i metodologii badania złączy *wet-mate* i *dry-mate*, używanych na dużych głębokościach w systemach podwodnych (ROV, sonary itp.).

Pośród specyfikacji złączy do aplikacji morskich w wielu miejscach przewijają się określenia „*dry-mate*” i „*wet-mate*”. Złącza tego pierwszego rodzaju to takie, które można łączyć i rozłączać tylko na sucho,

REKLAMA

PRODUCENT
**ELEMENTÓW
INDUKCYJNYCH**

www.feryster.pl

FERYSTER



Fotografia 14. Złącza firmy Harting w wykonaniu przeciwwybuchowym, spełniające wymagania norm ATEX (<https://connectorsupplier.com/>)

czyli poza środowiskiem wodnym, a dopiero potem zanurzać. Ich zaletą jest prostota i niezawodność w stałych instalacjach, gdzie kontakt z wodą następuje już po wykonaniu połączenia. Ich uszczelnienie jest pasywne i ma jedynie utrzymać szczelność zanurzonego połączenia.

Złącza typu *wet-mate* to specjalistyczne rozwiązania umożliwiające łączenie i rozłączanie również pod wodą, nawet w wymagających, morskich warunkach. Mogą być łączone nawet na dużych głębokościach. Zostały zaprojektowane tak, aby zapewniać bezpieczne i pewne połączenie elektryczne mimo obecności wody między stykami. Konstrukcja takiego złącza często stworzona jest tak, aby wypierać wodę z powierzchni styku podczas łączenia. Dzięki zaawansowanym mechanizmom uszczelniającym i zastosowaniu materiałów odpornych na korozję, takich jak wysokiej jakości elastomery oraz metale o podwyższonej odporności, można wykonywać operacje łączenia nawet przez zdalnie sterowane pojazdy podwodne lub nurków, bez ryzyka uszkodzenia komponentów. To rozwiązanie jest szczególnie przydatne w zastosowaniach, które wymagają częstej wymiany sprzętu, na przykład przy obsłudze robotów podwodnych, sensorów zatapialnych czy urządzeń na platformach wiertniczych.

W segmencie złączy przeznaczonych do zastosowań morskich (zarówno wiertniczych, jak i podwodnych) działa kilku wyspecjalizowanych producentów, którzy oferują odpowiednie certyfikaty i dopuszczenia. Do najważniejszych należą MacArtney (marka SubConn), produkujący kompletne złącza *wet-mate* i *dry-mate* używane w podwodnych pojazdach zdalnie sterowanych, sonarach i systemach oceanograficznych. TE Connectivity jest jednym z największych dostawców złączy hermetycznych spełniających wymagania IP68/IP69K do instalacji okrętowych i czujników głębokości. Firma Glenair oferuje złącza morskie oparte na konstrukcjach MIL-DTL-38999, ale wykonane z miedzioniklu i stali duplex, przeznaczone do stałego zanurzenia. Warto wymienić też firmy Hydro Group, SEALCON, Birns Aquamate oraz Amphenol Subsea, które dostarczają zarówno standardowe złącza do zastosowań podwodnych, jak i rozwiązania wykonywane pod konkretną specyfikację (tzw. *custom-made*), np. dla platform wiertniczych i systemów monitoringu oceanicznego.

Złącza dla przemysłu petrochemicznego i rafineryjnego

W środowisku petrochemicznym i rafineryjnym złącza muszą pracować w strefach zagrożonych wybuchem, gdzie obecność wybuchowych par lub pyłów stanowi wysokie zagrożenie dla

pracowników oraz samej infrastruktury. Z tego powodu producentów złączy do tych aplikacji obowiązują wymagania ATEX (w Europie) lub IECEx (na poziomie międzynarodowym) określające, że wszystkie komponenty elektryczne – w tym złącza – muszą mieć konstrukcję uniemożliwiającą zapłon atmosfery zewnętrznej. Jednym z kluczowych podejść jest tu iskrobezpieczeństwo, definiowane poprzez energię doprowadzoną do złącza oraz energię możliwą do zgromadzenia na stykach złącza. Musi ona być na tyle niska, żeby nawet podczas zwarcia lub uszkodzenia przewodów nie mogła powstać iskra o energii wystarczającej do zapłonu mieszaniny.

Iskrobezpieczeństwo osiąga się zarówno poprzez odpowiedni projekt elektryczny systemu (ograniczenie energii, elementy zabezpieczające, separacja torów), jak i poprzez dobór właściwych materiałów oraz konstrukcji złączy. Obudowy wykonywane są z materiałów nieiskrzących (np. stal nierdzewna, mosiądz niklowany, aluminium z powłoką anodowaną), a same styki mają zwiększony odstęp i poziom izolacji, aby uniemożliwić powstanie wyładowania łukowego. Bardzo ważne jest także hermetyczne uszczelnienie tylnej części złącza – każda szczelina mogłaby stać się kanałem zapłonu. Dlatego większość złączy ATEX/Ex ma specjalne uszczelki z elastomerów odpornych na paliwa oraz precyzyjnie dopasowane powierzchnie współpracujące (zdefiniowane przez normę jako „*spoiny zabezpieczające*”).

Tego typu konstrukcje muszą spełniać wymagania norm ATEX EN 60079-11 (złącza urządzeń iskrobezpiecznych) oraz IEC 60079-25 (instalacje i okablowanie w strefach zagrożenia wybuchem). Każde złącze przeznaczone do zastosowań iskrobezpiecznych jest certyfikowane według konkretnej klasy temperatury (np. T4 lub T5), która gwarantuje, że temperatura powierzchni nigdy nie przekroczy wartości mogącej zapalić pary znajdujące się w otoczeniu.

W wielu aplikacjach rafineryjnych stosuje się również złącza wyposażone w blokady mechaniczne, które uniemożliwiają przypadkowe rozłączenie lub utratę kontaktu pod obciążeniem. Najczęściej stosowane są blokady śrubowe lub obrotowe, które wymagają użycia narzędzia, zanim złącze zostanie odłączone – dzięki temu operator nie może przypadkowo doprowadzić do pojawienia się łuku przy rozłączeniu przewodu pod napięciem.

Z kolei w strefach o wysokim zapyleniu (instalacje przesypowe, silosy, pomieszczenia z pyłem koksowniczym) wykorzystuje się złącza z osłonami przeciwpyłowymi oraz kapturkami płomieniochronnymi. Ich zadaniem jest zapobieganie przedostawaniu się cząstek stałych do wnętrza korpusu oraz odseparowanie potencjalnej iskry od atmosfery zewnętrznej w taki sposób, aby nie mogło dojść do zapłonu. Ten typ rozwiązań certyfikuje się zgodnie z EN 60079-0 (ogólne wymagania urządzeń przeciwwybuchowych) oraz EN 60079-31, która określa dodatkowe wymagania dla osprzętu stosowanego w strefach zagrożenia wybuchem pyłów. W praktyce oznacza to nie tylko użycie odpowiednich materiałów uszczelniających, ale też ściśle zdefiniowanych wymiarów szczelin, promieni i głębokości prowadnic – wszystkie one muszą uniemożliwić wydostanie się płomienia lub rozgrzanych gazów na zewnątrz złącza, nawet w przypadku uszkodzenia wewnętrznego.



Fotografia 15. Popularne złącze trójfazowe 380 V w wykonaniu spełniającym normy ATEX (<https://rs-online.com/>)

Na rynku złączy w wykonaniu ATEX/Ex działa kilka wyspecjalizowanych firm, które oferują zarówno standardowe serie iskrobezpieczne, jak i złącza z blokadą mechaniczną oraz ochroną przeciwpyłową. Do najważniejszych producentów należą: Hawke International, Connomac, Amphenol Industrial oraz Radiall, R. STAHL.



Fotografia 16. Złącze RJ45 w wykonaniu ATEX
(<https://element14.com>)

Inne niszowe segmenty

Oprócz wymienionych powyżej branż istnieje jeszcze szereg innych, które stawiają złączom dość wysokie wymagania. Pomimo mniejszej skali zastosowań, wymagania środowiskowe oraz normy jakościowe bywają tu równie restrykcyjne.

W kolejnictwie podstawową rolę odgrywają normy EN 45545 (bezpieczeństwo pożarowe) oraz EN 50155 (urządzenia elektroniczne stosowane w taborze kolejowym). Złącza muszą sprostać bardzo surowym wymaganiom związanym z odpornością środowiskową oraz niezawodnością eksploatacji. Muszą wytrzymywać powtarzalne wstrząsy i wibracje występujące w czasie jazdy, pracować w szerokim zakresie temperatur (od ok. -55°C do +200°C), zapewniać skuteczne ekranowanie EMI/RFI oraz szczelność, jak również spełniać wymagania normy EN 45545 dotyczące ograniczenia ognia i dymu w wypadku pożaru. Omawiane złącza zwykle muszą mieć ponadto klasę szczelności IP66 lub IP67. W praktyce przekłada się to na zastosowanie trudnopalnych tworzyw na wkładki i korpusy, a także specjalnych powłok zabezpieczających styki przed korozją, wynikającą z długotrwałej ekspozycji na wilgoć, smary i zanieczyszczenia występujące w infrastrukturze kolejowej. Konektory te często są wyposażone w szybkozłącza bagnetowe i mechaniczne wzmocnienia, co gwarantuje pewne połączenie nawet w ekstremalnych warunkach.

Jeśli chodzi o producentów złączy kolejowych, na rynku wyróżniają się następujące firmy:

- ITT Veam (seria CIR) – oferuje złącza nie tylko elektryczne, ale i optyczne oraz pneumatyczne w wykonaniu hermetycznym, odpornym na paliwa i ekranującym przed EMI/RFI,
- Smiths Interconnect – serie LHS/LHZ i REP to modułowe złącza określane mianem *rackandpanel* o konstrukcji typu *blindmate*, o dużej gęstości styków, wysokiej odporności na drgania i uderzenia, a także wykonane z materiałów nisko-dymnych i ognioodpornych, zgodnych z normami kolejowymi,
- JAE – oferuje wysokoprądowe, uszczelniane złącza zgodne z FST,
- Binder – producent złączy w standardzie M12 z kodowaniem D, które stosowane są często np. w systemach kamer nadzorujących. Złącza te mają klasę IP67 i wyposażone są w ekrany elektromagnetyczne,
- Souriau (serie SMS) oferują złącza szybkiego podłączania, a także hermetyczne złącza BNC/TNC o wysokiej trwałości (do 500 połączeń) i szerokim zakresie temperatur pracy (od -65°C do 165°C), spełniające wymagania norm kolejowych.

W energetyce wiatrowej złącza narażone są na stały kontakt z wilgocią, bryzą morską (w przypadku instalacji typu *offshore*) oraz promieniowaniem UV. Komponenty używane w turbinach wiatrowych (zarówno na lądzie, jak i na morzu) muszą spełniać wymagania normy IEC 61400-1 oraz IEC 61400-4. Normy te określają m.in.: zakres temperatury pracy od -40°C do 85°C (dla komponentów wewnętrznych wieży) oraz od -40°C do 105°C (dla komponentów w gondoli i układzie mocy). Dodatkowo, złącza dla energetyki wiatrowej

muszą wykazywać odporność na wibracje i uderzenia mechaniczne, mierzone zgodnie z metodami podanymi w IEC 60068-2-6 (test wibracyjny) oraz IEC 60068-2-27 (test uderowy) i klasę IP co najmniej IP65 dla połączeń wewnętrznych gondoli oraz IP67/IP69K w przypadku połączeń zewnętrznych (zgodnie z IEC 60529).

W konstrukcjach morskich turbin wiatrowych wymagane jest dodatkowo potwierdzenie odporności na korozję, zwykle poprzez test w mgłę solną zgodnie z IEC 60068-2-11 lub ISO 9227. Złącza muszą również spełniać wymagania towarzystw klasyfikacyjnych (np. DNV-CG-0339 lub DNV-RU-OSS-1), które wymagają nie tylko odporności na mgłę solną, ale również na promieniowanie UV. W przypadku kabli/złączy dużej mocy (połączenie generator – transformator – rozdzielnia) stosuje się również normy IEC 60502-4 i IEC 61238. Dodatkowo wszystkie elementy instalowane na platformach offshore podlegają normom ATEX/IECEx (EN 60079-0, -7 oraz -31), jeżeli pracują w strefie potencjalnie zagrożonej atmosferą wybuchową (np. na platformach serwisowych lub przy układach hamowania turbiny).

Dominującymi na rynku producentami, którzy mają w swojej ofercie złącza z pełną kwalifikacją zgodną z IEC 61400 oraz IEC 60068, są firmy Harting, Amphenol Industrial, TE Connectivity, Weidmüller oraz LAPP. Dodatkowo w systemach wiatrowych, zwłaszcza w zastosowaniach związanych z integracją sensorów i systemów monitorowania wibracji w turbinie, stosuje się lekkie i uszczelnione złącza Binder, Smiths Interconnect (serie M12/M23) oraz Souriau (UTS), certyfikowane wg IEC 60529 do IP69K i zatwierdzone do stosowania w elektrowniach wiatrowych przez DNV-GL.

Trendy i przyszłość

W najbliższych latach można spodziewać się dalszego zacierania granic między poszczególnymi technologiami transmisji sygnału, czego najlepszym przykładem jest integracja światłowodów z przewodami miedzianymi w obrębie jednego złącza. Rozwiązania tego typu stają się szczególnie istotne w wymagających aplikacjach wojskowych i lotniczych, gdzie konieczne jest równoczesne przesyłanie sygnałów o dużej przepustowości, sygnałów sterujących oraz zasilania – przy zachowaniu pełnej hermetyczności i odporności na zakłócenia elektromagnetyczne. Firmy takie jak Axon' Cable, Solifos czy LEMO mają w swojej ofercie tego rodzaju złącza w wykonaniu wzmocnionym czy zgodnym z normami wojskowymi (**fotografia 18**).

Równolegle obserwujemy rosnącą popularność rozwiązań modułowych, które umożliwiają łatwe konfigurowanie złączy pod konkretne wymagania aplikacyjne, skracając czas projektowania i ułatwiając serwisowanie w terenie. Tego rodzaju złącza w jednej



Fotografia 17. Złącza CIR firmy VEAM spełniające wymagania kolejowe
(<https://connectorsupplier.com/>)



Fotografia 18. Złącze firmy Solifos integrujące w jednym konektorze połączenie elektryczne i optyczne (<https://solifos.com/>)

obudowie mogą mieć zainstalowany szereg różnych modułów, które dobiera się do konkretnej aplikacji. Przykładem może być seria Han-Modular firmy Harting, w których można łączyć różne moduły, w tym optoelektroniczne czy pneumatyczne(!). Złącza te dostępne są w wykonaniu wzmocnionym, jednakże nie oferują klasyfikacji wojskowej czy lotniczej.

Ważnym kierunkiem rozwoju pozostaje także miniaturyzacja – zarówno samych złączy, jak i akcesoriów – przy jednoczesnym utrzymaniu lub nawet podwyższeniu parametrów wytrzymałościowych. Dotyczy to szczególnie segmentów takich jak kosmonautyka czy lotnictwo, w których złącza pracują w ekstremalnych warunkach środowiskowych, ale istotne są także ich rozmiary i waga. Dążenie do zmniejszenia masy oraz poprawy odporności na korozję czy wytrzymałości na wysokie temperatury przekłada się natomiast na rozwój nowych materiałów kompozytowych (np. wypełnianych włóknem węglowym) oraz zastosowanie zaawansowanych powłok PVD, które pozwalają znacząco zwiększyć żywotność elementów, jednocześnie zapewniając lepsze parametry elektryczne i mechaniczne.



Fotografia 19. Modułowe złącze Han-Modular firmy Harting (<http://hartingconnector.com/>)

Podsumowanie

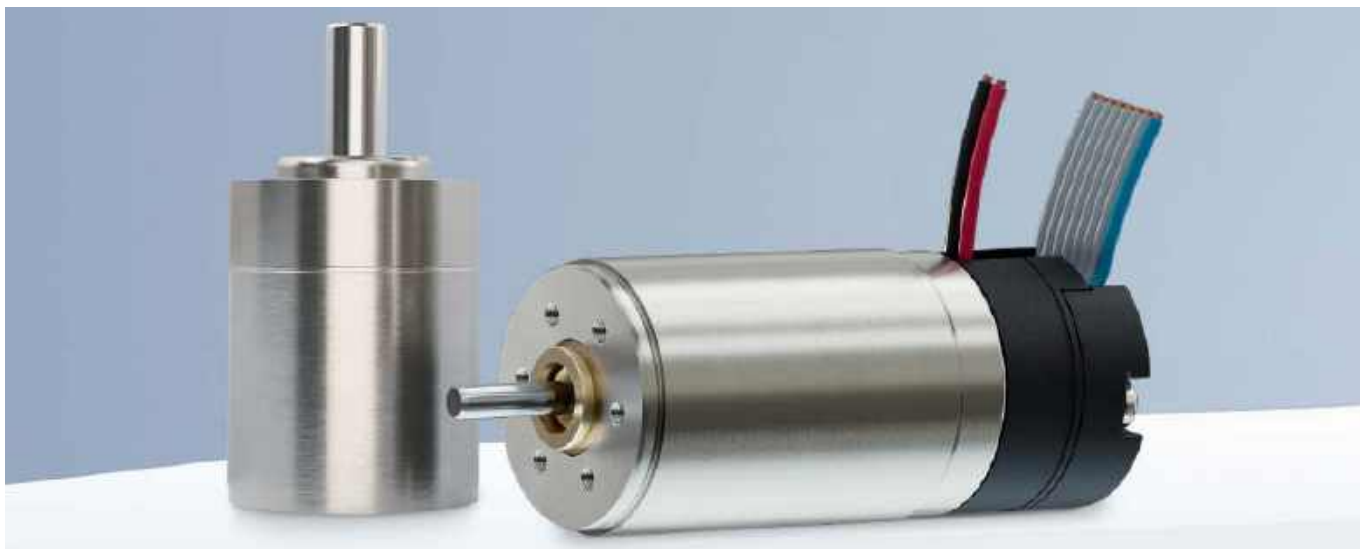
Złącza – choć często traktowane jako element drugoplanowy w całym projekcie systemu – stają się coraz bardziej zaawansowanymi komponentami, których rozwój będzie miał kluczowe znaczenie dla niezawodności nowoczesnej elektroniki. Jest to szczególnie widoczne w miejscach, gdzie stawiane złączom wymagania są wysokie.

Powyższy opis z pewnością nie wyczerpuje wszystkich rodzajów złączy, które przeznaczone są do pracy w wymagających warunkach. Intencjonalnie pominięto na przykład złącza do pracy w systemach medycznych, które wymagają ekstremalnie wysokiej niezawodności. Analogicznie pominięto rynek motoryzacyjny, gdyż jest on niezmiernie szeroki i sam opis złączy z odpowiednimi certyfikatami przekroczyłby objętość tego, i tak już dosyć długiego, artykułu.

Nikodem Czechowski, EP

Tabela 4. Podsumowanie wybranych obszarów aplikacyjnych złączy

Segment/ branża	Normy i standardy	Typowe wymagania środowiskowe	Materiały korpusów i styków	Przykłady ty- pów złączy	Przykładowi producenci
Wojskowa	MIL-DTL-38999, MIL-DTL-26482, MIL-DTL-5015, VG95234	Wibracje, wstrząsy, mgła solna, -65... 200°C, EMI/RFI	Aluminium z powłoką kadmową, stal nierdzewna, miedź, złoto	Okrągłe wielopinowe, hermetyczne, z filtrem EMI	Amphenol, TE Connectivity, Souriau, ITT Cannon
Lotnicza	ARINC, EN 2997, MIL-DTL-83723	Wysoka odporność na wibracje i zmiany temperatur, miniaturyzacja	Stopy aluminium, tytan, złożone styki	Micro-D, złącza awioniczne, światłowodowe	Glenair, Souriau, Radiall
Kosmiczna	ECSS-Q-ST-70, NASA-STD	Próżnia, promieniowanie, ekstremalne temperatury (>300°C), niskie odgazowanie	PEEK, PTFE, ceramika, złoto	Nano-D, hermetyczne, optyczne, do paneli słonecznych	Axon' Cable, TE Connectivity, Smiths Interconnect
Morska	DNV-GL, IEC 60092	Zasolenie, wilgoć, IP68/IP69K, korozja	Brąz, stal duplex, tytan, elastomery	Złącza podwodne wet/dry-mate	MacArtney, SubConn, SEACON
Petrochemia/ Rafinerie	ATEX, IECEx	Odporność chemiczna, iskrobezpieczeństwo, pyło- i gazoszczelność	Stal nierdzewna, aluminium anodowane	Wielopinowe, iskrobezpieczne, z blokadą	Connomac, Hawke, Bartec
Kolejnictwo	EN 50155, EN 45545	Odporność na wstrząsy, ognioodporność, wahania temperatury	Tworzywa trudnopalne, aluminium	Złącza prostokątne, modułowe	Harting, Stäubli, Schaltbau
Energetyka wiatrowa	IEC 61400	UV, wilgoć, korozja, zmiany temperatur	Kompozyty, stal nierdzewna	Złącza dużej mocy, hermetyczne	Lapp, Weidmüller, Huber+Suhner



Nowe silniki, przekładnie i enkodery o średnicy Ø 16 mm – doskonała synergia zapewniająca maksymalną wydajność

W ofercie FAULHABER pojawiły się nowe dodatki umożliwiające tworzenie kompletnych rozwiązań w zakresie technologii napędów. Dzięki nowym rozmiarom silników SXR, mocnemu silnikowi z nowej serii GXR, precyzyjnemu enkoderowi i kompatybilnej przekładni, projektanci mogą wybierać produkty idealnie pasujące do siebie, pochodzące z jednego źródła i zgodne ze standardową średnicą Ø 16 mm. Takie połączenie zapewnia optymalną wydajność, maksymalną dynamikę i absolutną precyzję – cechy te idealnie sprawdzają się w branżach opartych na zaawansowanych technologiach mechatronicznych, zwłaszcza w automatyce przemysłowej, robotyce i aparaturze medycznej.

Nowość w ofercie produktów: silniki DC z serii GXR i SXR

Nowy silnik szczotkowy 1627 GXR z miedziano-grafitowym układem komutacji imponuje mocą i szerokim zakresem opcji wyposażenia, dzięki którym może spełniać wymagania nowoczesnych rozwiązań napędowych. W ofercie dostępne są warianty na napięcie zasilania od 4,5 V do 24 V, z różnymi konfiguracjami łożysk. Istnieje również możliwość indywidualnego dostosowania silnika – począwszy od modyfikacji przedniego i tylnego wału, aż po opcje pozwalające na pracę w próżni lub w wysokiej temperaturze otoczenia. Zoptymalizowane wyważenie wirnika zapewnia płynną pracę i zwiększa trwałość napędu. Technologia uzwojenia z sześciokątnymi cewkami o wysokim współczynniku wypełnienia miedzią i zoptymalizowaną prostą sekcją, a także wysokiej jakości magnesy zapewniają stabilność termiczną i poprawiają wydajność silnika.

W serii SXR z układem komutacji ze stali nierdzewnej również pojawiła się nowość – do istniejących modeli 1218 i 1228 SXR dołączyła nowa wersja w rozmiarze 1627 SXR. Charakteryzuje się optymalnym stosunkiem mocy do wielkości i doskonale nadaje się

Więcej informacji:

FAULHABER Polska sp. z o.o.
ul. Kosynierów 44, Sosnowiec
tel. +48 61 278 72 53
faulhaber.com/en/, info@faulhaber.pl



do zastosowań w branży zaawansowanych technologii. Wszystkie komponenty serii SXR i GXR są zgodne z dyrektywą RoHS. Warto też dodać, że do dyspozycji projektantów jest szereg opcji konfiguracyjnych w zakresie połączeń elektrycznych.

Silniki GXR i SXR można łatwo łączyć z metalowymi przekładniami planetarnymi z serii GPT. Nowy model kompatybilny ze standardem 16 mm (16GPT) doskonale sprawdzi się w przypadku wymagających zastosowań o ograniczonej przestrzeni montażowej. Zoptymalizowana konstrukcja pozwala na osiągnięcie dużych prędkości i wykorzystanie pełni możliwości silnika. Ponadto stabilna konstrukcja gwarantuje niezawodne przenoszenie ekstremalnych sił i bezproblemową pracę przy dużych obciążeniach.

Kolejny sojusznik: nowy enkoder magnetyczny IEX3

Dzięki najnowocześniejszym czujnikom scalonym modele IEX3 i IEX3 L oferują wysoką rozdzielczość i precyzję pozycjonowania z dokładnością do 0,3°. Szerokie zakresy napięcia zasilania (3,3 V do systemów akumulatorowych oraz 5 V – do większości pozostałych zastosowań) i temperatury (od -40 do 100°C) sprawiają, że enkoder zapewnia elastyczność aplikacyjną i odporność na trudne warunki pracy. Model IEX3 (L) jest dostępny w wersjach ze sterownikiem liniowym lub bez niego, oferuje ponadto niezwykle kompaktowe rozmiary i łatwość konserwacji – to idealne rozwiązanie do stosowania w połączeniu z nowymi silnikami FAULHABER SXR i GXR.

Dzięki idealnemu dopasowaniu komponentów projektanci i inżynierowie mogą korzystać z kompaktowego, wydajnego i niezawodnego kompletnego rozwiązania, które otwiera nowe możliwości przed konstruktorami nowoczesnych systemów napędowych.



Odpadowe metody obróbki metalu

Metale to – obok izolatorów i półprzewodników – najczęściej stosowane materiały w elektronice. Doskonale przewodzą prąd i ciepło, mają szeroki zakres właściwości magnetycznych i wykazują wysoką odporność na zmiany temperatury. Są również bardziej wytrzymałe mechanicznie niż pozostałe materiały używane w urządzeniach elektronicznych. Wszystkie te cechy pozwalają na tworzenie z metalu trwałych i wytrzymałych elementów mechanicznych, takich jak zębatki przekładni, sprężła, elementy nośne, obudowy oraz wszelkiego rodzaju złącza i styki. Ze względu na liczne funkcje i zastosowania metali ważne jest, aby znać podstawowe metody ich obróbki. W naszym poprzednim artykule (EP 07/2025) opisaliśmy silnie rozwijające się w ostatnich latach, bezodpadowe technologie wytwarzania z metalu. Teraz postaramy się przybliżyć bardziej klasyczne (odpadowe) metody obróbki, które pozwalają uzyskać pożądany kształt i właściwości powierzchni detali przez usuwanie nadmiaru materiału. Niektóre metody, które uznaliśmy za drogie lub trudno dostępne, umyślnie pominęliśmy.

Metale w elektronice są używane przede wszystkim do wytwarzania płytek drukowanych (warstwa miedzi), wszelkiego rodzaju styków lub elementów mechanicznych (głównie stal, mosiądz i aluminium), a także obudów (aluminium lub stal), ekranów urządzeń (stopy żelaza, miedź i proszki miedziane) oraz wielu innych elementów. Warto zwrócić uwagę na fakt, że produkcja części metalowych metodą odpadową (stosowaną głównie w przypadku komponentów mechanicznych czy niektórych obudów) jest z zasady

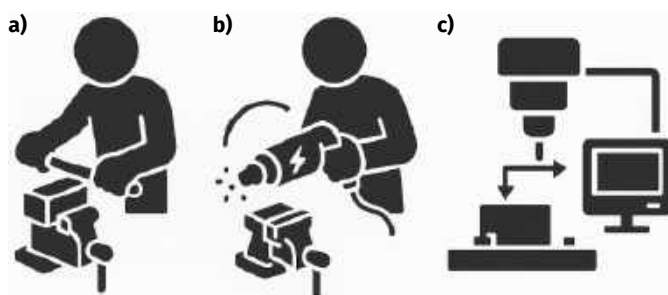
bardziej skomplikowana (pod względem liczby operacji, które należy wykonać) od metod przyrostowych bądź bezodpadowych.

Jako przykład weźmy panel czołowy urządzenia. Gdy wydrukujemy go na drukarce 3D (np. metodą SLM) jedyne, co będziemy musieli zrobić, by panel był funkcjonalny, to nadrukować odpowiednie opisy. W przypadku obróbki odpadowej sprawa prezentuje się nieco inaczej. Przede wszystkim musimy z kawałka litego metalu, najlepiej w kształcie płyty, wyciąć na określony wymiar jej fragment, w którym – usuwając materiał odpowiednimi narzędziami – ukształtujemy geometrię zewnętrzną oraz wszystkie niezbędne otwory, a czasem także gwinty. Następnym krokiem jest wykończenie powierzchni i krawędzi poprzez odpowiednie szlifowanie lub polerowanie, a dopiero po zakończeniu tych procesów przystąpić do malowania i nadruku oznaczeń na obudowie. Podsumowując można stwierdzić, że technologia odpadowa jest dużo bardziej skomplikowana od technologii przyrostowej i bardzo często wieloetapowa. Niemniej jednak niższy koszt wykonania detalu metodą obróbki odpadowej, a także dostępność maszyn i materiałów, mogą zachęcać do korzystania z tej tradycyjnej technologii. Także właściwości mechaniczne tak wytworzonych detali są zwykle lepsze niż w przypadku wykonania ich techniką przyrostową.

Niniejszy krótki przegląd metod obróbki odpadowej stanowi systematyczny opis głównych właściwości, wad i zalet tych technologii w kontekście zastosowań w elektronice. Uwaga będzie skupiona na głównych rodzajach obróbki maszynowej, z pominięciem najprostszych metod obróbki ręcznej (brzeszczotem, pilnikiem, prostymi elektronarzędziami etc.) i bardziej zaawansowanych, hybrydowych technik obróbki maszynowej.

Obróbka ręczna i maszynowa

Zanim omówimy ważniejsze metody obróbki odpadowej, podzielone według kryterium rodzaju użytych narzędzi oraz ruchów narzędzi i obrabianego detalu, warto zauważyć, że większość z nich może być



Rysunek 1. Podział metod obróbki na przykładzie szlifowania: technika ręczna – pilnikiem (a), ręczna – elektronarzędziem szlifierskim (b) i maszynowa CNC (c)

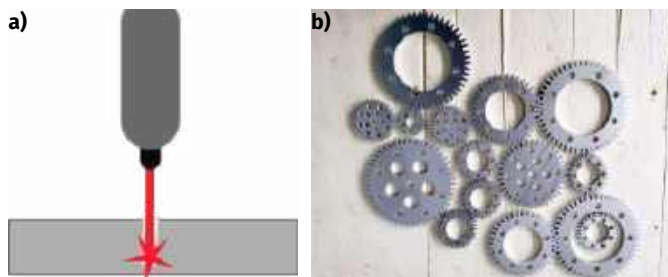
zrealizowana ręcznie lub maszynowo. W obróbce ręcznej narzędzie (np. brzeszczot, pilnik, nóż) lub elektronarzędzie (np. piła tarczowa, szlifierka, wiertarka) obsługiwane są przy pomocy siły mięśni człowieka. Z kolei w obróbce maszynowej ruch narzędzia i obrabianego przedmiotu są realizowane z wykorzystaniem napędu mechanicznego. Klasyczne obrabiarki to stacjonarne elektronarzędzia wyposażone w manualne pokręta lub koła, które połączone są z osiami sterującymi narzędziem lub stołem obróbkowym. Rozwój elektroniki i informatyki spowodował, że tradycyjne maszyny obróbkowe, oparte na mechanicznym lub elektromechanicznym sterowaniu, są zastępowane przez maszyny z komputerowym sterowaniem numerycznym (ang. *Computer Numerical Control*, skrót CNC). Proces obróbki z wykorzystaniem maszyn CNC składa się z trzech etapów:

- projektowania CAD (przygotowanie modelu 3D elementu w programie komputerowym),
- programowania CAM (model 3D jest przetwarzany na kod sterujący maszyną),
- właściwej obróbki (maszyna wykonuje działania zgodnie z zaprogramowanymi instrukcjami).

Zaletą stosowania maszyn CNC jest obróbka z dużą szybkością, dokładnością i powtarzalnością oraz ograniczenie błędów ludzkich i strat materiałowych. Schematyczną ilustrację obróbki szlifowania realizowanego ręcznie, przy pomocy elektronarzędzia i maszyny CNC zaprezentowano na **rysunku 1**.

Cięcie

Metody cięcia metali można podzielić na mechaniczne, termiczne, erozyjne i elektrochemiczne. Cięcie mechaniczne korzysta z narzędzi tnących, takich jak: piły tarczowe, piły taśmowe czy nożyce gilotynowe i służy do cięcia po linii prostej. Cięcie mechaniczne wykrojnikami pozwala uzyskać skomplikowane kształty z blach. Z kolei cięcie termiczne (gazowe, laserowe lub plazmowe) polega na topieniu i usuwaniu metalu w miejscu zaplanowanej linii podziału. Metody erozyjne to głównie cięcie wodne lub elektroerozyjne, w których – odpowiednio – strumień wody lub wyładowania elektryczne usuwają materiał w celu jego podzielenia. Podobny cel w metodzie elektrochemicznej realizują reakcje chemiczne. Poniżej omówimy bardziej zaawansowane technicznie, a zarazem dokładniejsze metody cięcia.



Rysunek 2. Schematyczna ilustracja procesu cięcia laserowego (a) i przykłady elementów mechanicznych wyciętych tą metodą [1] (b)

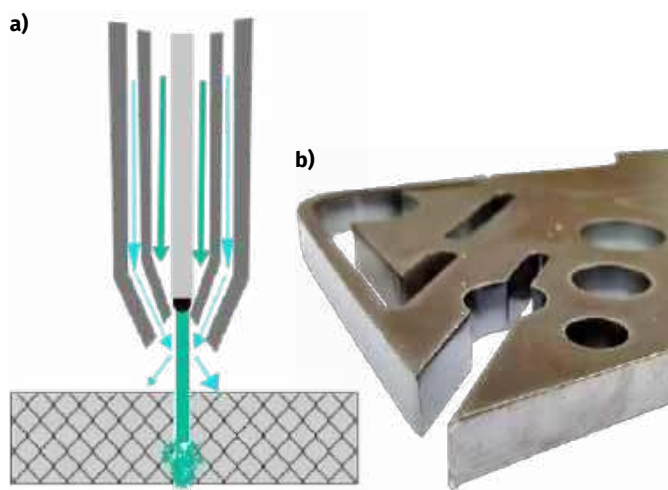
Cięcie laserowe

Technologia cięcia laserowego metalu polega na stopniowym topieniu lub odparowaniu metalu za pomocą skupionej wiązki laserowej i odbywa się – w zależności od technologii – warstwa po warstwie lub od razu na całej grubości blachy. Cięcie laserem oferuje wysoką dokładność, choć nie tak wysoką, jak frezowanie czy toczenie. Typowa dokładność cięcia laserowego to 0,1 mm. Ograniczeniem opisywanej techniki jest trudność w cięciu metali nieżelaznych (takich jak miedź czy aluminium), ze względu na ich wysoki współczynnik odbicia światła.

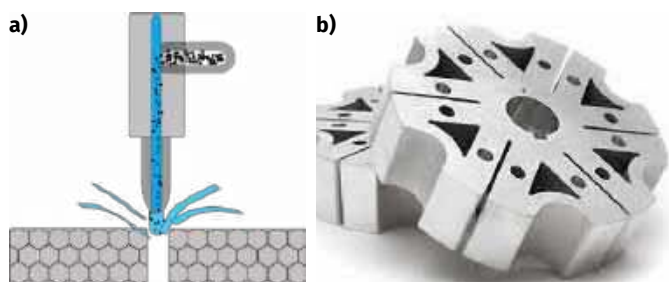
Lasery można sklasyfikować na lasery MOPA (ang. *Master Oscillator Power Amplifier*) ze statyczną głowicą oraz lasery z głowicą ruchomą (plotery laserowe). Te ostatnie są wolniejsze i nadają się głównie do cięcia, natomiast lasery MOPA dobrze sprawdzają się w grawerowaniu. Wśród technologii laserowych najpopularniejsze odmiany to: światłowodowe (ang. *fiber laser*), CO₂, UV oraz IR. Ceny laserowych urządzeń do cięcia lub grawerowania zaczynają się już od 5000 zł. Alternatywą dla kompletnego urządzenia do laserowego cięcia jest kupno modułu laserowego do frezarki/drukarki 3D, o ile dane urządzenie może obsługiwać takie dodatkowe wyposażenie. Przy wyborze lasera warto zwrócić uwagę, by jego moc wynosiła co najmniej 80...160 W. Metodą cięcia laserowego można z dużą dokładnością wykonywać np. szablony do nakładania pasty lutowniczej, a nawet same płytki drukowane, jak również wszelkiego rodzaju obudowy składane i zginane z blach. Należy pamiętać, że cięcie blach o znacznej grubości budżetowym laserem jest trudne i może wymagać dużo czasu.

Cięcie plazmą

Cięcie plazmą odbywa się przez topienie metalu przez strumień zjonizowanego, gorącego gazu (plazmy). Przecinarki plazmowe średniej klasy oferują dokładność cięcia metalu wynoszącą od 0,2 mm do 0,5 mm, zatem wybór tej metody może nie być najlepszym rozwiązaniem, gdy wymagana jest duża dokładność. W przypadku takich prac, jak przygotowanie materiału do wykonania paneli przednich czy obudów lub ekranów na poszczególne części urządzenia (czujniki, układy radiowe) cięcie plazmą jest jednak optymalnym rozwiązaniem ze względu na koszt i dostępność technologii, w porównaniu do cięcia laserowego czy wodnego. Cięcie plazmą jest dość tanią technologią i, gdy dysponujemy frezarką/ploterem, możemy za około 5000 zł dokupić do niej moduł cięcia plazmowego. Gdy planujemy zaopatrzyć się w kompletną maszynę, musimy liczyć się z cenami na poziomie od około 10 000 zł.



Rysunek 3. Schematyczna ilustracja procesu cięcia plazmowego (a) i przykład elementu wyciętego tą metodą [2] (b)



Rysunek 4. Schematyczna ilustracja procesu cięcia wodą (a) i przykład elementu wyciętego tą metodą [3] (b)

Cięcie wodą (ang. *waterjet*)

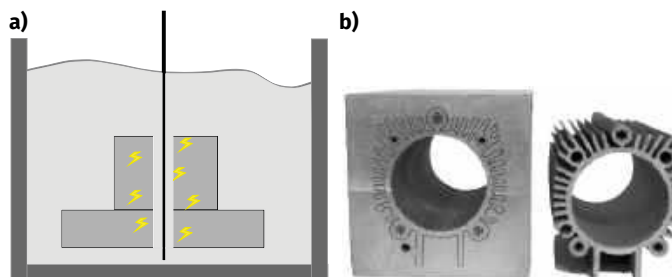
Technologia waterjet polega na wykorzystaniu strumienia wody lub wody z ziarnami piasku (bądź innych kryształów) pod wysokim ciśnieniem, nawet do kilkudziesięciu tysięcy barów. Dysza, z której wytryskuje woda, ma małą średnicę, co pozwala na precyzyjne i niskoubytkowe cięcie. Wodą można obrabiać bardzo grube blachy, często o grubości ponad 10 mm. Klasyczne maszyny do cięcia wodą są duże i kosztowne, a ponadto mogą nie zmieścić się w małym warsztacie. Do zalet technologii waterjet należą wysoka szybkość i powtarzalność, zaś minusami tej technologii są: wysoki koszt oraz znaczne zużycie energii, wynoszące zazwyczaj od 10 do 15 kW. W ostatnim czasie pojawiły się na rynku maszyny typu desktop, które można zakupić w cenie od 30 000 zł.

Cięcie elektro-drutowe (WEDM – *Wire Electrical Discharge Machining*)

Ciekawą technologią cięcia metali jest elektrodrążenie drutowe, które opiera się na wyładowaniach elektrycznych powodujących lokalne odparowanie metalu. Cała operacja odbywa się w wodzie dejonizowanej lub oleju – obrabiany przedmiot podłączony jest do masy, natomiast cienki drut obrabiający (elektroda) – do generatora impulsów. Dzięki zastosowaniu bardzo cienkiego drutu obrabiającego technologia ta charakteryzuje się niską ilością odpadów. Wadą technologii WEDM jest niemożność drążenia otworów – w ten sposób można jedynie powiększać istniejące otwory. Jest to metoda dokładna, jednak jeszcze niezbyt powszechna na rynku w postaci rozwiązań typu desktop. Warto wiedzieć, że firma Rack Robotics oferuje swoje przecinarki WEDM w przedsprzedaży za cenę od 8 do 12 000 zł. Inne maszyny typu desktop można okazjonalnie znaleźć w cenie od 10 000 do 25 000 zł. Z kolei przemysłowe rozwiązania tego typu są znacznie droższe i wykraczają poza możliwości finansowe mniejszych firm.

Skrawanie

Obróbkę skrawaniem można podzielić na: wiórową oraz ścierną. Obróbka wiórowa to np. toczenie, frezowanie, wiercenie, struganie, dłutowanie czy gwintowanie – materiał usuwa się za pomocą ostrzy narzędzia skrawającego. Z kolei obróbka ścierna to szlifowanie, honowanie, docieranie lub polerowanie – usuwanie materiału następuje za pomocą ścierniwa typu papier lub płótno ścierne,



Rysunek 5. Schematyczna ilustracja procesu cięcia metodą elektro-drutową (a) i przykład elementów wyciętych tą metodą [4] (b)



Rysunek 6. Przykłady maszyn do cięcia blach: programowalna maszyna do cięcia laserowego [5] (a) i ręcznie obsługiwana przecinarka plazmowa [6] (b)

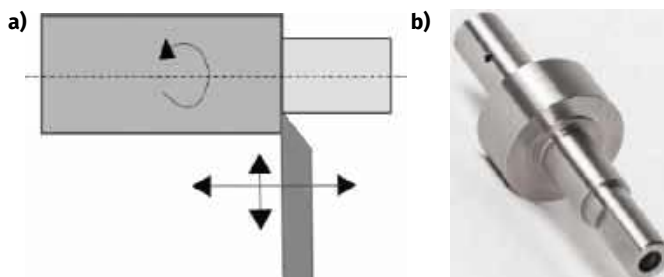
ściernice, proszki czy pasty. Obróbka wiórowa powoduje usuwanie większych porcji materiału, natomiast podczas szlifowania niwelowane są drobne nierówności z powierzchni lub krawędzi, jest to zatem obróbka płytka. Szlifowanie stanowi więc zwykle następny etap po obróbce wiórowej. Poniżej omówimy ważniejsze oraz stale rozwijane techniki skrawania.

Toczenie

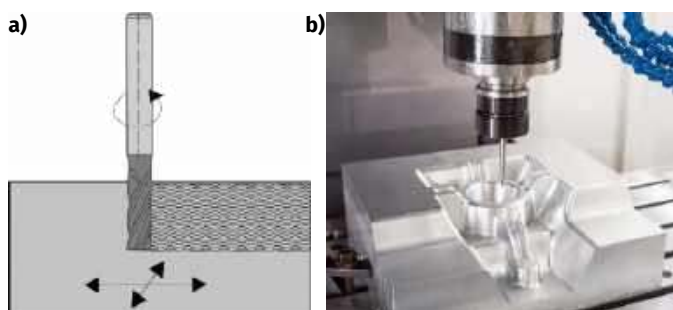
Toczenie to rodzaj obróbki, w której obracany jest zamocowany w uchwycie przedmiot, a narzędzie skrawające wykonuje ruch w płaszczyźnie zawierającej oś obrotu przedmiotu. Technika toczenia przydaje się szczególnie, gdy chcemy uzyskać osie lub wałki napędowe, tuleje lub pierścienie oraz inne kształty o symetrii osiowej. Do urządzeń elektronicznych można toczyć na przykład cylindryczne obudowy lub elementy mechanizmów, sensorów, części do silników, enkoderów lub potencjometrów. Dokładność tokarek CNC to zazwyczaj 0,01 mm, zaś przy użyciu tokarek manualnych zależy ona od umiejętności operatora, a także od poziomu drgań urządzenia. Tokarka może być użyta do wiercenia otworów w osi elementów zamocowanych w uchwycie (wiertło mocowane jest wówczas w tulei konika) lub też do gwintowania czy szlifowania (tarcza szlifierska jest w takim przypadku mocowana w specjalnej przystawce lub w uchwycie tokarki). Ceny tokarek manualnych do metalu zaczynają się od około 3500 zł, zaś tokarka CNC to trochę większa inwestycja, od 25 000 zł za nową maszynę. Do ceny tokarki należy doliczyć także koszty niezbędnych narzędzi, czyli noży tokarskich i osprzętu.

Frezowanie

Frezowanie jest obróbką wirującym narzędziem skrawającym (frezem), podczas gdy obrabiany przedmiot pozostaje



Rysunek 7. Schematyczna ilustracja procesu toczenia (a) i przykład wytoczonego elementu [7] (b)

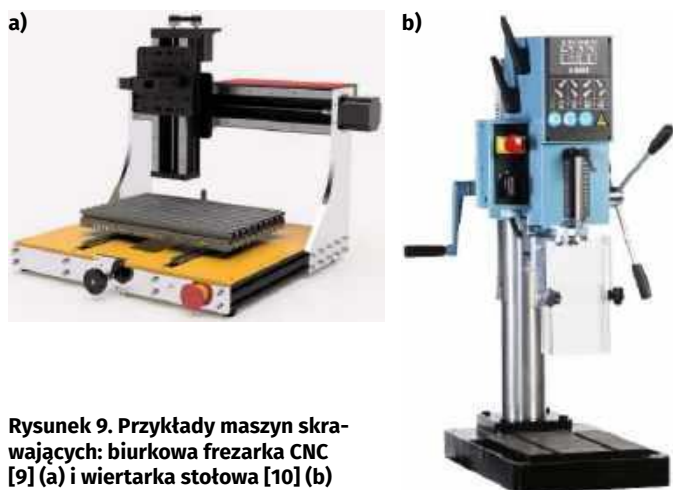


Rysunek 8. Schematyczna ilustracja procesu frezowania (a) i fotografia obróbki frezowaniem [8] (b)

przymocowany do precyzyjnie kontrolowanego, ruchomego stołu. Technologia ta pozwala na wykonanie detali o skomplikowanych kształtach, nadaje się też do precyzyjnego wiercenia otworów. W przypadku frezarek CNC proces obróbki skrawaniem może być kontrolowany nie tylko w trzech osiach, ale nawet w większej ich liczbie. Obróbka frezarką CNC przydaje się w warsztacie elektronika, bowiem metodą tą można wykonywać panele przednie, całe obudowy z metalu, przeróżne elementy mechaniczne (takie jak zębatki czy zaciski), ale także jedno- lub dwuwarstwowe płytki PCB. Kluczem do doboru optymalnej frezarki jest zdefiniowanie potrzeb. Większość frezarek CNC oferuje dokładność 0,01 mm. Przy frezowaniu należy zwrócić uwagę na poprawne zamocowanie obrabianego elementu oraz odpowiednie chłodzenie. Frezowanie nie jest łatwą techniką, lecz przy odrobinie wysiłku można wykonywać bardzo dobrej jakości części do prototypów urządzeń. Najtańsze frezarki manualne są dostępne w cenie od około 1000 zł, a koszty budżetowej maszyny CNC do metalu zaczynają się od około 1500 zł na chińskich platformach sprzedażowych lub 2500 zł w europejskich sklepach.

Wiercenie

Wiercenie to proces obróbki skrawaniem, którego celem jest wykonywanie otworów w materiale. Obok właściwego wiercenia wyróżnia się obróbki pokrewne, takie jak powiercanie lub rozwiercanie (powiększanie lub zwiększanie dokładności wykonania otworu), czy nawiercanie lub pogłębianie (wykonanie niewielkiego otworu przydatnego w dalszym wierceniu lub zwiększanie średnicy otworu do pewnej głębokości, np. jako gniazdo dla śruby). Technika wiercenia jest bardzo często wykorzystywana w pracy elektronika, m.in. do wykonania połączeń mechanicznych i elektrycznych (zacisków), otworów w panelach czołowych i obudowach, płytkach PCB, a także otworów pod gwinty wewnętrzne itp. Wiercenie wykonuje się przede wszystkim za pomocą przenośnych lub stacjonarnych wiertarek, ale można je także realizować



Rysunek 9. Przykłady maszyn skrawających: biurkowa frezarka CNC [9] (a) i wiertarka stołowa [10] (b)



Rysunek 10. Przykłady gwintowników i narzynek [11] [12] [13] (a) oraz gwinciarko-wiertarka stołowa [14] (b)

na frezarkach (mocując wiertła zamiast frezów) lub tokarkach. Ten ostatni sposób dotyczy głównie elementów walcowych, w których należy wykonać otwór precyzyjnie w osi detalu. Ceny prostych wiertarek stołowych rozpoczynają się od około 500 zł, jednakże koszt wiertarek wyposażonych w osprzęt pozwalający na precyzyjne wiercenie, zwłaszcza w przypadku maszyn CNC, może być znacznie wyższy.

Gwintowanie

Gwintowanie to obróbka niezbędna w projektach, w których przewiduje się połączenia śrubowe bez zastosowania nakrętek. Z wyjątkiem toczenia, które pozwala na wykonanie gwintów zewnętrznych, inne omawiane techniki obróbki nie dają możliwości wykonania gwintów. Tu z pomocą przychodzą różne rodzaje manualnych gwintowników (do wykonania gwintów wewnętrznych) i narzynek (do wykonania gwintów zewnętrznych), a także specjalne maszyny do gwintowania. Narzędzia stosowane przez maszyny gwintujące różnią się od narzędzi stosowanych przez maszyny gwintujące. Niektóre z gwintów w aparaturze elektronicznej wymagają zastosowania specjalnych rozwiązań lub narzędzi, które gwarantują wykonanie bardzo dokładnych gwintów, na przykład w celu mocowania radiatorów i innych komponentów, czy też w procesie produkcji obudów wodoszczelnych i pyłoszczelnych. Takie specjalne gwinty zapewniają stabilne i szczelne połączenie elementów, bez ryzyka luzowania w wyniku drgań. Ceny maszyn do gwintowania (znanych także pod nazwą gwinciarek) zaczynają się od około 2000 zł. Alternatywą dla gwintowania może być montaż w otworach, np. za pomocą kleju, gotowych wkładek gwintowanych – dotyczy to jednak głównie elementów z tworzyw polimerowych.

Szlifowanie i obróbka końcowa

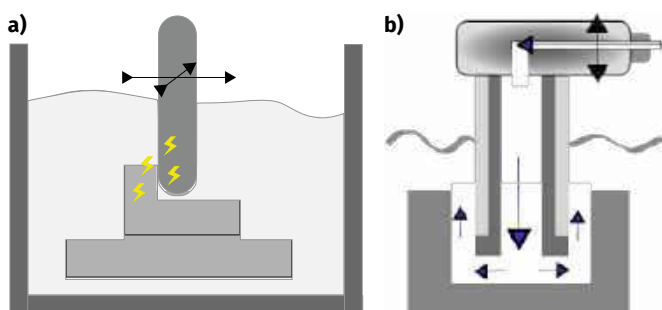
Szlifowanie jest obróbką wykańczającą, stosowaną na przykład gdy toczony lub frezowany przedmiot ma nabrać estetycznych walorów użytkowych, bez widocznych śladów jakiejkolwiek wcześniejszej obróbki zgrubnej. Innym powodem stosowania szlifowania mogą być względy bezpieczeństwa (usuwanie ostrych krawędzi) czy konieczność dopasowania łączonych lub współpracujących



Rysunek 11. Przykłady urządzeń do szlifowania: dwutarcowa szlifierka stołowa [15] (a), miniszlifierka manualna [16] (b) oraz polerka bębnowa [17] (c)

ze sobą elementów. Klasyczne, ręczne szlifowanie odbywa się za pomocą pilników lub papieru ściernego i jest to proces znany zapewne wszystkim Czytelnikom. Jednak gdy chcemy, by szlifowanie przebiegało sprawniej i dokładniej, powinniśmy użyć do tego celu sprzętu elektromechanicznego, takiego jak szlifierki stołowe czy taśmowe, które oferują szybszą i bardziej jednorodną obróbkę elementów. Jeszcze większe możliwości dają szlifierki CNC, zapewniające wyższą jakość i powtarzalność procesu wygładzania powierzchni w porównaniu do narzędzi ręcznych lub elektronarzędzi. Pozwalają ponadto na zachowanie dużej dokładności wymiarowej i skrajnie wąskich tolerancji (poniżej 0,02 mm). Kolejnym po szlifowaniu etapem obróbczym może być polerowanie. Jest to często długi proces, trwający nawet wiele godzin i przydatny szczególnie wtedy, gdy nie zamierzamy malować, lakierować lub anodować obrabianego elementu. Dodatkową zaletą polerowania jest fakt, że zapewnia ono możliwość osiągnięcia niskiego współczynnika tarcia. Obok samej maszyny szlifierskiej bardzo ważny w tego rodzaju obróbce jest właściwy dobór ścierniwa (pasty lub proszku).

Zakup najprostszej szlifierki stołowej to wydatek około 200 zł, natomiast ceny prostych szlifierek taśmowych lub bębnowych oscylują wokół 300 zł. Wadą takich przyrządów jest niemożność



Rysunek 12. Schematyczna ilustracja obróbki elektroerozyjnej (a) i elektrochemicznej (b)

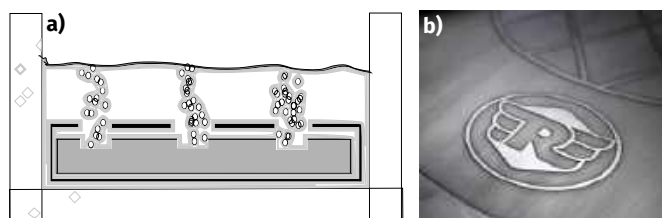
szlifowania bardziej złożonych geometrycznie detali. Ciekawą opcją, głównie przydatną w prototypowaniu, są miniszlifierki, pozwalające na dokonywanie niewielkich zmian geometrii detalu lub usuwanie niedokładności/błędów produkcyjnych. Porównywalnym pod względem cenowym urządzeniem, lecz przeznaczonym do większych poprawek lub wiercenia otworów nieuwzględnionych w projekcie, jest minifrezarka. Tego typu elektronarzędzie to wydatek około 150 zł.

Obróbka elektroerozyjna EDM (Electrical Discharge Machining) i elektrochemiczna

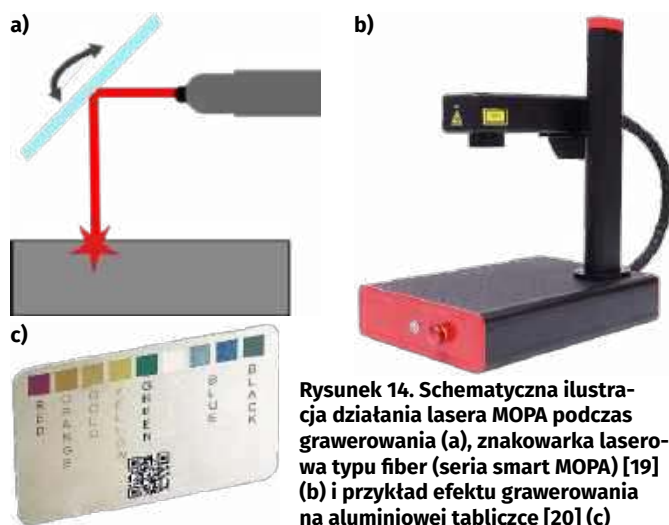
Obróbka elektroerozyjna polega na drążeniu metalu w wyniku działania wyładowań elektrycznych pomiędzy elektrodą a obrabianym detalem, jednocześnie zanurzonymi w nieprzewodzącym płynie (dejonizowanej wodzie lub oleju). Energia cieplna generowana przez wyładowania elektryczne powoduje topienie lub parowanie materiału z przedmiotu obrabianego. Kontrola geometrii drążonego obszaru odbywa się poprzez odpowiedni dobór kształtu elektrody oraz sterowanie jej ruchem.

Technologia EDM nie jest w stanie wykonywać ostrych krawędzi oraz wycinać precyzyjnie w blachach. Z uwagi na wymienione ograniczenia często oferowane są systemy hybrydowe, łączące klasyczne EDM z drutowym WEDM. Takie rozwiązanie pozwala wykonywać precyzyjne, wręcz „lustrzane” elementy i detale, które mogą być docelowo wykorzystane w aplikacjach wymagających niewielkiego tarcia. Ceny ręcznych maszyn EDM do drążenia otworów zaczynają się od 1000 zł, natomiast koszt maszyny CNC to już około 100 000 zł – można jednak skorzystać z oferowanych w dość przystępnej cenie usług firm posiadających tego typu maszyny.

Obróbka elektrochemiczna, podobnie jak elektroerozyjna, także bazuje na przepływie prądu, jednak proces erozji nie zachodzi pod wpływem wytworzonego ciepła, ale wskutek reakcji chemicznych. Pomiędzy przedmiotem obrabianym (anodą) i narzędziem (katodą) znajduje się przepływający przez cały czas obróbki strumień elektrolitu. Stały prąd jonowy pomiędzy elektrodami i towarzyszące mu reakcje chemiczne powodują usuwanie materiału z obrabianego detalu i wypłukiwanie go przez przepływający elektrolit. Wielkość i kształt erodowanej objętości jest odwzorowaniem zastosowanej katody. W odróżnieniu od obróbki elektroerozyjnej, opisywana metoda nie wiąże się z generowaniem napiężeń cieplnych



Rysunek 13. Schematyczna ilustracja trawienia chemicznego płytki PCB (a) i przykład wytrawionego wzoru [18] (b)



Rysunek 14. Schematyczna ilustracja działania lasera MOPA podczas grawerowania (a), znakowarka laserowa typu fiber (seria smart MOPA) [19] (b) i przykład efektu grawerowania na aluminiowej tabliczce [20] (c)

w obrabianym detalu, bo proces ten nie powoduje znaczącego podgrzewania elementu.

Trawienie chemiczne

Metoda polega na usuwaniu cienkich warstw metalu przez wykorzystanie reakcji chemicznych w roztworach kwasów (np. azotowego, fluorowodorowego czy siarkowego) w odpowiednich stężeniach, czasem wspomaganych dodatkiem soli. W wyniku trawienia usuwana jest z powierzchni niepożądana warstwa metalu lub jego związków. Proces ten jest stosowany do oczyszczania (np. z tlenków, zgorzelin czy rdzy) i przygotowania powierzchni do dalszej obróbki (np. malowania, powlekania, spawania) oraz do wytwarzania precyzyjnych komponentów lub obróbki dekoracyjnej. Technologia ta nadaje się do wytwarzania precyzyjnych obudów dla czujników czy małych elementów mechanicznych, stanowi ponadto podstawową metodę wytwarzania płytek PCB. Podczas takiego procesu na laminat pokryty warstwą miedzi najpierw nanosi się warstwę ochronną w miejscach, w których mają pozostać ścieżki i wylewki. Warstwa ochronna może być wykonana poprzez nałożenie materiału światłoczułego (fotorezystu), selektywnie utwardzonego za pomocą fotolitografii. Obszary miedzi niepokryte warstwą ochronną usuwane są właśnie poprzez trawienie, zaś utwardzony fotorezyst ze ścieżek jest następnie czyszczony rozpuszczalnikiem. Cena procesu (w najprostszym wydaniu) zazwyczaj nie przekracza znacząco kosztu wytrawiacza oraz wanny na płyny, należy natomiast uwzględnić koszty nadrukowania lub nałożenia warstwy ochronnej.

Grawerowanie

Celem grawerowania jest naniesienie na detal napisów lub wzorów, głównie poprzez płytką obróbkę ubytkową technikami mechanicznymi bądź laserowymi. Grawerować możemy w metalach, a także w praktycznie każdym innym materiale, niemal niezależnie od kształtu. Stosując technikę laserową mamy możliwość dynamicznego ustalenia energii, co pozwala na wytwarzanie warstwy tlenków o różnych, ustalonych programowo barwach. Grawerowanie na panelach urządzeń elektronicznych pozwala na nadawanie im walorów estetycznych oraz użytkowych, na przykład przez wykonanie symbolu marki czy opisów przycisków. Opisywana metoda pozwala także na modyfikowanie powierzchni w ściśle określonych celach technicznych – graweruje się na przykład stoły do drukarek 3D, aby zwiększyć przyczepność. W porównaniu do druku na elementach metalowych czy nanoszenia naklejek z nadrukiem, grawerowanie jest trwalsze i pozwala zachować pożądane funkcje w warunkach, w których nadruk czy naklejka uległyby uszkodzeniu. Ceny grawerek mechanicznych zaczynają się od około 1000 zł,



podczas gdy niedrogi ploter laserowy CNC można kupić za około 600 zł. Warto zauważyć, że istnieje opcja doposażenia grawerek oraz ploterów o tzw. czwartą oś, czyli dodatkowy, kontrolowany numerycznie uchwyt obrotowy, który pozwala na grawerowanie na elementach walcowych.

Podsumowanie

Z uwagi na różnorodność odpadowych technik wytwarzania elementów metalowych i ograniczoną objętość niniejszego artykułu, wiele dostępnych metod zostało pominiętych. Nie wspomnieliśmy też o możliwościach wytwarzania elementów metalowych przy użyciu technik trwałego łączenia zamiast usuwania materiału (takich jak klejenie, spawanie czy nitowanie elementów metalowych). Warto podkreślić, że znajomość technik obróbki metali a także ich zalet i wad, może pozwolić na optymalny wybór metody wytwarzania, zarówno z punktu widzenia jakości produktu, jak też ceny procesu. Wiedza ta może być także przydatna przy wyborze wyposażenia naszego warsztatu oraz podejmowaniu decyzji, które prace należy zlecić firmom zewnętrznym, a które możemy wykonać sami. Należy jednak pamiętać, by przed zaplanowaniem rozpoczęcia obróbki dokonać przeszukania w bazach internetowych lub serwisach aukcyjnych – być może da się znaleźć gotowe elementy, które potrzebujemy. Zakup takowych pozwoli zaoszczędzić czas i często obniżyć koszt wytworzenia urządzenia.

**Stanisław Kaczmarek
prof. Mariusz Kaczmarek¹**

¹ Wydział Mechatroniki, UKW w Bydgoszczy

Źródła fotografii

- [1] <https://t.ly/Cx4Gn>
- [2] <https://t.ly/SiMHP>
- [3] <https://t.ly/FuKY>
- [4] https://t.ly/qd_kC
- [5] <https://t.ly/8pyRT>
- [6] <https://t.ly-sVks>
- [7] <https://t.ly/o707C>
- [8] <https://t.ly/iGO83>
- [9] <https://t.ly/tiF-8>
- [10] <https://t.ly/iGXb0>

- [11] <https://t.ly/ys1Xk>
- [12] <https://t.ly/AJb6J>
- [13] <https://t.ly/NaWiy>
- [14] <https://t.ly/27gYs>
- [15] <https://t.ly/WJCNT>
- [16] <https://t.ly-ToPl>
- [17] <https://t.ly/JZo0->
- [18] <https://t.ly/89zMg>
- [19] <https://t.ly/5DFVC>
- [20] <https://t.ly/ZBIIL>



Proste sposoby na napięcie ujemne

We współczesnej elektronice na ogół używamy napięć zasilających o wartości dodatniej względem potencjału masy. Napięcia ujemne towarzyszą najczęściej układom analogowym i niektórym wyświetlaczom. Jak można takie napięcie uzyskać możliwie prostym sposobem?

W artykule prezentuję kilka sprawdzonych sposobów.

Poniedziałek, 9:00, telefon od klienta.

– Halo, Michał?

– Nooo, cześć, co Cię sprowadza?

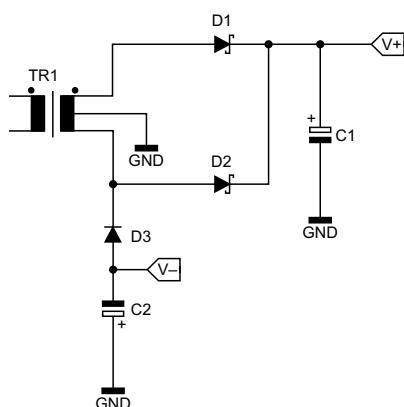
– Słuchaj, pamiętasz ten ostatni prototyp zasilacza, który niedawno nam wysłałeś?

– Pamiętam, wszystko z nim w porządku?

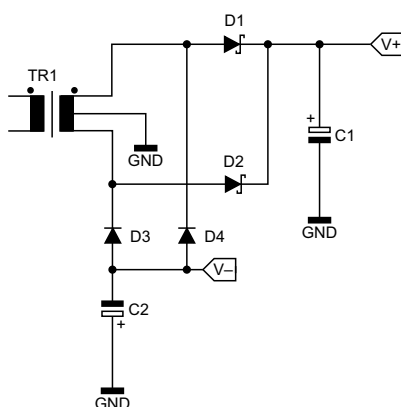
– W porządku jak najlepszym, tylko zapomnieliśmy, że do jednego opampa przyda się kawałek minusa, tak chociaż z -4 V . Wydajność prądowa bez znaczenia. Ogarniesz coś?

– Coś się wymyśli – odeślij, dolutuję to i owo.

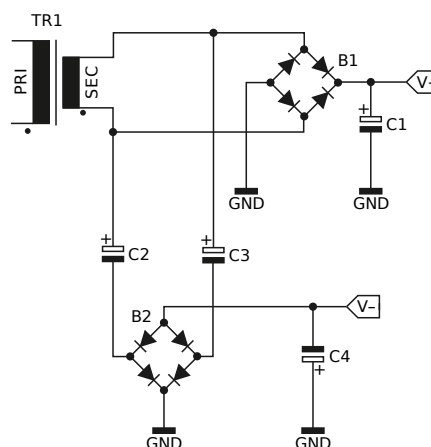
Jeżeli układ z zamierzenia ma korzystać z napięć ujemnych, na przykład do zasilania wzmacniaczy operacyjnych lub innych układów analogowych, to uwzględnia się ten fakt na etapie projektowania zasilacza do owego układu. Zostaje wtedy użyty transformator z dwoma uzwojeniami, drugi prostownik, drugi filtr, stabilizatory etc. Problem pojawia się, kiedy mamy układ zasilany napięciem dodatnim (złożony np. z mikrokontrolera, wyświetlacza, modułu Wi-Fi), a wspomniane napięcie ujemne jest potrzebne tylko dla



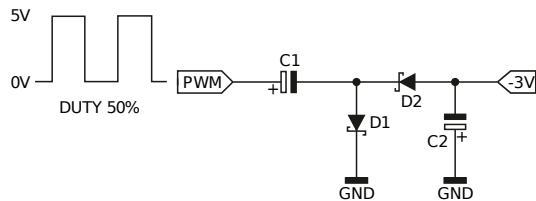
Rysunek 1. Prostownik jednopółkowy z transformatorem o podzielonym uzwojeniu



Rysunek 2. Uzupełnienie układu z rysunku 1 do postaci prostownika dwupółkowego



Rysunek 3. Rozwiązanie stosowane w przypadku klasycznego zasilacza z mostkiem Graetza



Rysunek 4. Prosty układ „odwracający” napięcie zasilany impulsowo

jednego, małego wzmacniacza operacyjnego, który wzmacnia napięcie z jakiegoś czujnika. Chodzi tylko o to, aby ów wzmacniacz prawidłowo przetwarzał sygnał bliski 0 V, więc to ujemne napięcie nie musi być „kosmicznie” wysokie – z reguły wystarczy potencjał względem masy na poziomie $-4...-3$ V. Wymagana wydajność prądowa też jest w takich sytuacjach niewysoka, wynosi przeważnie kilkanaście miliamperów. Mało tego – nawet stabilizacja nie jest konieczna, gdyż PSRR wzmacniacza i tak robi swoje.

I co począć w takim przypadku? Dostawiać drugi transformator albo przetwornicę izolowaną do wytworzenia tego napięcia? Niby można, ale w wielu przypadkach wcale nie trzeba tego robić. Przecież to są dodatkowe koszty, masa, rozmiary – oszczędzajmy wszędzie tam, gdzie ma to sens.

W przypadku zasilacza omawianego na wstępie mieliśmy do czynienia z transformatorem sieciowym o dzielnym uzwojeniu, który dostarczał napięcie dodatnie. W prostowniku znalazły się diody Schottky'ego (D1 i D2), co pozwoliło zminimalizować spadek napięcia i zapewnić możliwe wysokie napięcie wejściowe stabilizatora liniowego. W tej sytuacji dodałem tylko dwa elementy, jak na **rysunku 1**. Dioda D3 „wyciąga” prąd z kondensatora C2 w ujemnych półokresach sinusoidy, w których dioda D2 jest zatkana. Rozwiązanie bardzo proste, tanie i zapewniało napięcie ujemne o niezbyt wygórowanych parametrach. W tym układzie to wystarczyło.

Z racji niesymetrycznego obciążenia uzwojeń tworzy się składowa stała prądu w uzwojeniu wtórnym, co podmagnesowuje rdzeń. Trzeba zwrócić na to uwagę przy pobieraniu większego prądu z zasilacza napięcia ujemnego, bo w takich warunkach rdzeń traci przenikalność, co pogarsza sprawność. Wystarczy tylko jedna modyfikacja w postaci dodania diody D4 (**rysunek 2**). Można zauważyć, że w tej chwili cały prostownik tworzy elegancki mostek Graetza.

Jednak nie zawsze mamy do dyspozycji transformator z dzielnym uzwojeniem wtórnym. Prawdę mówiąc, w nowoczesnych układach to dosyć rzadko spotykane rozwiązanie. Częściej mamy

do czynienia z sytuacją, w której pojedyncze uzwojenie wtórne jest podłączone do wejścia prostownika w układzie Graetza. Czymś w pełni prawidłowym byłoby użycie drugiego uzwojenia, do zasilania drugiego prostownika – jednak nie zawsze jest to możliwe i/lub opłacalne. Użycie układu jak z rysunku 1 lub rysunku 2 nie jest możliwe, coś więc zrobić? Można „dobudować” prosty obwód składający się z dwóch kondensatorów elektrolitycznych, które zapewnią polaryzację diod w drugim mostku Graetza (**rysunku 3**). W praktycznych zastosowaniach kondensatory C2 i C3 mają pojemność rzędu 1000 μ F lub więcej, co pozwala na pobranie z „ujemnego” wyjścia tak powstałego zasilacza prądu o natężeniu kilkudziesięciu miliamperów. Jediną przeszkodą jest tutaj reakcja kondensatorów. Trzeba również mieć na uwadze spadek napięcia na tychże kondensatorach, wywołany ową reaktancją.

No dobrze, a jeżeli w ogóle nie mamy do dyspozycji napięcia przemienne, to co wtedy? Mamy układ zasilany, dajmy na to, z akumulatora bądź zespolonej przetwornicy impulsowej. Układy typu ICL7660, dedykowane do wytwarzania napięcia ujemnego, są dosyć drogie i mają mocno ograniczone parametry. Wtedy pomocny może okazać się dobrze znany obwód z **rysunku 4**. Na jego wejście podaje się przebieg prostokątny o wypełnieniu 50% i częstotliwości rzędu kilkuset herców do kilku kiloherców. Im wyższa będzie jego amplituda, tym wyższe (co do wartości bezwzględnej) napięcie ujemne otrzymamy na wyjściu – w praktycznych realizacjach, przy użyciu diod Schottky'ego, trzeba liczyć się ze stratą napięcia wynoszącą około 2 V. Polecam przy tym stosować kondensatory elektrolityczne lub tantalowe o pojemności 100 μ F lub zbliżonej, choć ten układ jest w stanie wiele wybaczyć.

Z oczywistych względów należy zadbać o wysoką wydajność prądową wyjścia generującego ów przebieg. W swojej praktyce używam samodzielnego oscylatora na bazie układu 555 lub – jeżeli mam do tego warunki – korzystam z jednego z wbudowanych liczników mikrokontrolera. Wówczas ów licznik wytwarza ciągłą falę prostokątną całkowicie sprzętowo, bez angażowania rdzenia. Do zwiększenia amplitudy oraz wydajności prądowej polecam użycie scalonych driverów MOSFET. Bardzo dobrze nadają się one do tej roli – można je sterować niskim napięciem, a zasilac – wysokim (na poziomie nawet kilkunastu woltów), ponadto oferują małą rezystancję wyjściową. Jeden z moich znajomych porównał ten układ do powielacza napięcia, który bardzo chętnie był stosowany przez inżynierów z byłej NRD w urządzeniach RTV. I coś w tym jest...

Michał Kurzela, 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

Internet Rzeczy w pomiarach środowiskowych (21)

Pomiar pH

Monitorowanie jakości wody jest niełatwym zadaniem, a szczególnie dotyczy to pomiaru pH. Zadanie jest nieco ułatwione przy zastosowaniu odpowiedniej sondy pH z gotowym wzmacniaczem. Zaprezentowany system pomiaru pH i temperatury roztworów na bazie zestawu Gravity Lab Grade Analog pH Sensor Kit oraz płytki Raspberry Pico 2 w prosty sposób umożliwia spełnienie szeregu wymagań funkcjonalnych stawianych tego typu systemom.

Poziom pH roztworu jest jednym z często branych pod uwagę pomiarów w wielu branżach, np. rolnictwie czy medycynie. Taką funkcję realizują też urządzenia testowo-pomiarowe do laboratoriów, automatyki domowej (monitoring basenów), kontroli infrastruktury sieciowej (np. bieżące monitorowanie jakości wody pitnej) itd. Wiele gałęzi przemysłu wykorzystuje poziom pH do różnych celów. Wartość pH służy jako punkt odniesienia do monitorowania i zapobiegania korozji rur i kotłów. W procesach fermentacji ciągle pomiary pH są istotne, aby uniknąć wytwarzania niepożądanych i szkodliwych produktów ubocznych. W przemyśle piwowarskim pH służy do określania starzenia, wzrostu twardości chmielu i stężenia goryczki. W przypadku produktów łatwo psujących się, takich jak mięso i ryby, poziom pH określa okres przydatności do spożycia i świeżość.

System pomiaru pH musi spełniać szereg wymagań funkcjonalnych, takich jak:

- wysoka impedancja wejściowa do próbkowania napięcia elektrody pH,
- pomiar temperatury w celu kompensacji zależności pH od temperatury,
- niski pobór mocy, aby zmaksymalizować czas pracy na baterii,
- efektywne tłumienie szumów, aby zminimalizować niepewność pomiaru.

Skrót pH pochodzi od łacińskiego terminu „potentia hydrogenii” (ang. potential of hydrogen), co można przetłumaczyć jako „potencjał wodoru”. Skala pH to ilościowa skala kwasowości i zasadowości roztworów wodnych związków chemicznych, oparta na aktywności jonów wodorowych H^+ w roztworach wodnych. Dla celów pomiarowych norma ISO oraz Międzynarodowa Unia Chemii Czystej i Stosowanej (IUPAC) definiują kwasowość/zasadowość jako wartość pH roztworu, w którym jest zanurzone standardowe ogniwo galwaniczne

Substancja	pH
1 M kwas solny	0
0,1 M kwas solny	1
Sok żółtkowy	1,5 – 2
Sok cytrynowy	2,4
Coca-Cola	2,5
Oceł	2,9
Sok pomarańczowy	3,0
Piwo	4,5
Kawa	5,0
Herbata	5,5
Kwaśny deszcz	< 5,6
Mleko	6,5
Chemicznie czysta woda	7
Ślina człowieka	6,5 – 7,4
Krew	7,35 – 7,45
Woda morska	8,0
Mylidło	9,0 – 10,0
Woda amoniakalna	11,5
Wodorotlenek wapnia	12,5
1 M roztwór NaOH	14

Rysunek 1. Przykładowa reprezentacja skali pH [4]

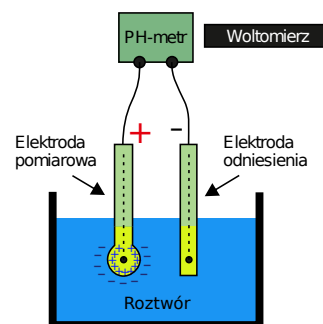


Poprzednie odcinki znajdują się pod adresem:
<https://ulubionykiosk.pl/media>

zdefiniowane przez IUPAC i dla którego zmierzono wartość pierwszej siły elektromotorycznej [5]. Z definicji tej wynika, że pH roztworów jest jednostką bezwymiarową i ma charakter jedynie porównawczy, nieprzekładający się bezpośrednio na stężenie czy aktywność jonów wodorowych ani żadnych innych. Definicja ta jest np. wykorzystywana przy przygotowywaniu skal dla papierków uniwersalnych oraz pH-metrów. Skala pH została zdefiniowana pierwotnie dla rozcieńczonych roztworów kwasów, zasad i soli, więc jej zastosowanie poza zakresem od 0 do 14 jest rzadko spotykane [5]. pH równe 7 jest uważane za neutralne, poniżej 7 – kwaśne, a powyżej tego progu – zasadowe (alkaliczne). Orientacyjnie skala pH została pokazana na **rysunku 1** [4].

Do pomiarów pH używa się zwykle papierków nasączonych mieszaniną substancji wskaźnikowych, które zmieniają kolor w szerokim zakresie. Dokładniejszych pomiarów dokonuje się metodą nazywaną pH-metrią. Większość pH-metrów to w istocie mierniki potencjału, w których pH ustala się na podstawie pomiaru siły elektromotorycznej (SEM) ogniwa utworzonego z elektrody wskaźnikowej (zanurzonej w roztworze badanym) i elektrody porównawczej (**rysunek 2**). Bardziej złożone pH-metry są dodatkowo wyposażone w termometry, gdyż temperatura ma wpływ na pomiar.

Istnieje wiele różnych konstrukcji elektrod pH-metrów. Najbardziej rozpowszechnione są jednak mierniki ze zintegrowanymi elektrodami w jednej sondzie, składającej się ze szklanej elektrody pomiarowej



Rysunek 2. Zasada działania pH-metru [4]

i elektrody odniesienia. Typowa konstrukcja bazuje na cienkiej, szklanej membranie, która otacza roztwór chlorowodoru (HCl). Wewnątrz obudowy znajduje się srebrny drut pokryty AgCl, który działa jako elektroda odniesienia i styka się z roztworem HCl. Jony wodoru, znajdujące się na zewnątrz szklanej membrany, dyfundują przez nią i wypierają odpowiednią liczbę jonów sodu (Na⁺), które normalnie występują w większości szkła. Jony dodatnie ograniczają się głównie do powierzchni szkła po tej stronie membrany, po której stężenie jest niższe. Nadmiar ładunku Na⁺ generuje napięcie na wyjściu czujnika.

Sonda jest analogiczna do baterii. Po umieszczeniu czujnika w roztworze elektroda pomiarowa generuje potencjał zależny od aktywności wodoru w roztworze, który jest porównywany z potencjałem elektrody odniesienia. Wraz ze wzrostem kwasowości roztworu (niższe pH) potencjał elektrody szklanej staje się bardziej dodatni (+mV) w porównaniu z elektrodą odniesienia, a wraz ze wzrostem zasadowości roztworu (wyższa wartość pH) potencjał elektrody szklanej staje się bardziej ujemny (−mV) w porównaniu z elektrodą odniesienia. Typowa sonda pH generuje idealnie 59,154 mV/jednostkę pH w temperaturze 25°C. Zwykle wyraża się to równaniem Nernsta [14].

Elektroda pH jest czujnikiem pasywnym, co oznacza, że nie wymaga źródła wzbudzenia (napięcia ani prądu). Ponieważ wyjście elektrody może wahać się powyżej i poniżej punktu odniesienia, jest ona klasyfikowana jako czujnik bipolarny. Generuje ona napięcie wyjściowe, które jest liniowo zależne od pH mierzonego roztworu.

Impedancja elektrody pH jest bardzo wysoka, ponieważ cienka szklana bańka ma dużą rezystancję, zazwyczaj mieszczącą się w zakresie od 10 do 1000 MΩ. Oznacza to, że elektrodę można monitorować wyłącznie za pomocą urządzenia pomiarowego o wysokiej impedancji wejściowej [15].

W roztworze stężenie jonów wodorowych może zmieniać się wraz z temperaturą. W szczególności wartość pH wzrasta w roztworach kwaśnych, wraz ze wzrostem temperatury, natomiast w roztworach obojętnych lub zasadowych zależność ta jest odwrotna. Aby zapewnić dokładne odczyty, niezbędna jest zatem korekta termiczna wartości pH w odniesieniu do wyniku pomiaru zarejestrowanego w temperaturze 25°C.

Biorąc pod uwagę wysoką impedancję wyjściową, najbardziej klasycznym podejściem jest zaprojektowanie aktywnego obwodu wejściowego o wysokiej impedancji, aby uniknąć przesunięcia odczytów napięcia, wynikającego z prawa Ohma. Konsekwencją obecności w urządzeniu obwodu o wysokiej impedancji wejściowej jest wysoki poziom szumów, obserwowany np. w środowiskach przemysłowych (często wyższy niż deklarowana dokładność systemu).

Istnieją dwa praktyczne rozwiązania powyższego problemu:

- filtrowanie w domenie analogowej (w dowolny sposób niezbędny do zniwelowania szumu w torze sygnałowym przed przetwornikiem ADC),
- filtrowanie w domenie cyfrowej (umożliwiające uzyskanie średniej ruchomej i odchylenia standardowego, z dodatkową korzyścią: można łatwo zidentyfikować i odrzucić wartości odstające od normy).

Zestaw Gravity: Lab Grade Analog pH Sensor Kit

Zestaw Gravity: Lab Grade Analog pH Sensor Kit for Arduino/Raspberry Pi (with Calibration Solutions) (SEN0161-V2) [1] jest przeznaczony do testowania pH w warunkach laboratoryjnych. Całość składa się z laboratoryjnej sondy pomiarowej pH (fotografia 1), płytki Gravity pH Meter V2.0 (fotografia 2) oraz buteleczek z roztworami wzorcowymi dla pH 4,0 oraz 7,0. Pojemniki są szczelnie zamknięte aluminiową folią, znajdującą się pod nakrętką.

Sonda pomiarowa pH [1]:

- zakres detekcji pH: 0...14,
- zakres temperatur: 5...60°C,
- punkt zerowy: $7 \pm 0,5$;
- powtarzalność: $< 0,017$,



Fotografia 1. Zestaw pomiarowy z sondą pH oraz roztworami wzorcowymi [1]



Fotografia 2. Płytkę pomiarową Gravity pH Meter V2.0 [1]

- szum: $< 0,5$ mV,
- czas reakcji: < 2 min,
- rezystancja wewnętrzna: < 250 MΩ,
- żywotność sondy: $> 0,5$ roku (w zależności od częstotliwości użytkowania),
- długość kabla: 100 cm.

Kabel sondy jest zakończony

wtyczką BNC służącą do podłączania do płytki pomiarowej.

Moduł pomiarowy V2 [1]:

- napięcie zasilania: 3,3...5,5 V
- napięcie wyjściowe: 0...3,0 V,
- złącze sondy: BNC,
- złącze sygnałowe: PH2.0-3P (Gravity),
- dokładność pomiaru: $\pm 0,1$ @ 25°C,
- wymiary: 42 mm $\times$ 32 mm/1,66 $\times$ 1,26 cala.

Na płytce pomiarowej został zastosowany precyzyjny wzmacniacz operacyjny LMC6482 typu CMOS z wejściem i wyjściem typu Rail-to-Rail oraz bardzo małym prądem wejściowym 0,02 pA (rezystancja wejściowa 10 TΩ). Sygnał pomiarowy jest sprzętowo filtrowany, co zapewnia niski poziom szumu.

Zalecenia dotyczące stosowania zestawu pomiarowego [2]:

1. Złącze BNC i płytkę konwersji sygnału muszą być suche i czyste, w przeciwnym razie wpłynie to na impedancję wejściową, co spowoduje niedokładny pomiar.
2. Płytkę konwersji sygnału nie może być umieszczona bezpośrednio na mokrej lub przewodzącej powierzchni. Zaleca się użycie nylonowego słupka do zamocowania płytki konwersji sygnału i zapewnienie pewnej odległości między płytką konwersji sygnału a powierzchnią podłoża.
3. Czuła szklana bańka w głowicy sondy pH nie powinna dotykać twardego materiału. Jakiegokolwiek uszkodzenia lub zarysowania spowodują uszkodzenie elektrody.
4. Po zakończeniu pomiaru należy odłączyć sondę pH od płytki pomiarowej. Sonda nie powinna być podłączana do płytki pomiarowej bez zasilania przez dłuższy czas.
5. Nakrętka butelki sondy zawiera płyn ochronny (3,3 mol/l KCL). Nawet jeśli nakrętka butelki jest mocno dokręcona, część płynu ochronnego może nadal wyciekać wokół nakrętki, tworząc białe kryształy. Jednak dopóki w nakrętce znajduje się płyn ochronny, nie wpłynie to na żywotność ani dokładność sondy. Zaleca się, aby białe kryształy umieścić z powrotem w płynie ochronnym w nakrętce butelki.
6. Po zakończeniu pomiaru należy założyć nasadkę ochronną i wlać do niej niewielką ilość roztworu KCL o stężeniu 3 mol/l, aby utrzymać szklaną bańkę wilgotną.

7. Gdy sonda jest używana po raz pierwszy lub nieużywana przez dłuższy czas, należy ją zanurzyć w roztworze 3N KCL na 8 godzin.
8. Po długotrwałym użytkowaniu należy sondę przemyć wodą destylowaną, a następnie zanurzyć w roztworze 3N KCL.
9. Należy unikać długotrwałego zanurzenia sondy w wodzie destylowanej, białku lub roztworze fluorku kwasu oraz unikać kontaktu z olejem silikonowym.
10. Przed pomiarem innego roztworu należy umyć sondę w wodzie destylowanej i zebrać resztki wody papierem, aby zapobiec zanieczyszczeniu krzyżowemu między roztworami.

Sondy pH podlegają dryfowi po dłuższych pomiarach. W przypadku konieczności pomiaru roztworu przez dłuższy czas, zalecane jest użycie zestawu *Gravity: 7/24 Industrial Analog pH Meter Kit (Arduino, Raspberry Pi)* (SEN0169-V2) [6].

Zmodyfikowany moduł DFRobot I²C ADS1115

Moduł DFRobot I²C ADS1115 (DFR0553) firmy DFRobot [8] zawiera układ przetwornika analogowo-cyfrowego ADS1115. W celu zapewnienia optymalnych warunków zasilania, płytkę należy zmodyfikować zgodnie z opisem, który zamieściliśmy w poprzednich odcinkach niniejszego cyklu. Do pracy z układem ADS115 została zastosowana biblioteka języka MicroPython [9].

Płytki RPi Pico 2 firmy Raspberry Pi

Nowe płytki Pico 2 i Pico 2W firmy Raspberry Pi z procesorem RP2350 są zgodne elektrycznie z płytkami Pico pierwszej serii (Pico/Pico W) [10]. Na płytkach zostały zastosowane układy pamięci NOR Flash serii W25Q (Winbond) o częstotliwości pracy do 133 MHz (przepustowość do 66 MB/s). Dokładny opis zamieszczono w artykule „Płytki Raspberry Pi Pico 2/2W z procesorem RP2350” [11].

Płytki Pico 2 zawierają przetwornicę buck-boost, która dostarcza napięcie 3,3 V (do zasilania RP2350 i obwodów zewnętrznych) z szerokiego zakresu napięć wejściowych (od 1,8 do 5,5 V). Zapewnia to znaczną elastyczność w zasilaniu urządzenia z różnych źródeł, takich jak pojedyncze ogniwo litowo-jonowe lub 3 ogniwa AA połączone szeregowo. Najprostszym sposobem zasilania Pico 2 jest podłączenie kabla do umieszczonego na płycie gniazda Micro USB. W dokumentacji płytki pokazano, jak poprzez dodanie tranzystora MOS można zrealizować podtrzymanie baterijne zasilania modułu [12].

Czujnik temperatury DS18B20

Pomiary standardowe i kalibracyjne napięcia dostarczanego przez sondę są uzupełniane przez jednoczesny pomiar temperatury przez dwa czujniki DS18B20. Do pomiaru temperatury płynów zastosowano sondy wodoodporne DFR0198 (IP68, obudowa ze stali nierdzewnej) z czujnikiem temperatury DS18B20 [7]. Czujniki działają w trybie 3-przewodowym. Pracują na szynie 1-Wire dołączonej do wyprowadzenia GP16 procesora, ze wspólnym rezystorem podciągającym 5 kΩ.

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 px, przeznaczony do płytek z serii Raspberry Pi Pico. Ma wbudowany kontroler, który realizuje komunikację za pomocą interfejsu SPI. Pico Graphics to zunifikowana biblioteka grafiki i wyświetlania firmy Pimoroni umożliwiająca sterowanie wyświetlaczami z Pico w języku MicroPython [16].

Ekspander szyny Pico

Ekspandery firmy Pimoroni są przeznaczone do płytek z serii Raspberry Pi Pico. Wyposażone zostały w jedno standardowe złącze żeńskie do bezpośredniego wpięcia płytki RPi Pico oraz zestawy męskich listew 2×20 pinów, które umożliwiają podłączenie

dotychczasowych modułów rozszerzeń. Etykiety wyprowadzeń, umieszczone na górnej stronie płytki, znacznie ułatwiają prototypowanie. Ekspander Pico Decker (Quad Expander) (PIM555) ma cztery zestawy męskich listew, a Pico Omnibus (Dual Expander) PIM556 – dwa zestawy.

Przygotowanie środowiska programowego

Interpreter MicroPython firmy Pimoroni dla RPi Pico 2 zawiera dodatkowo sterowniki wielu czujników oraz wyświetlaczy, w tym Pico Inky Pack [17].

1. Zmontuj elementy zgodnie z opisem.
2. Pobierz najnowszy interpreter MicroPythona w pliku *Pico2-v0.0.12-pimoroni-micropython.uf2* ze strony firmy Pimoroni [17].
3. Trzymając wciśnięty biały przycisk BOOTSEL, podłącz Raspberry Pi Pico 2 do komputera kablem microUSB.
4. Skopiuj pobrany plik *.uf2* do pamięci Raspberry Pi Pico 2. Jest ono widoczne jako dysk RP2350 w eksploratorze plików Windows.
5. W komputerze zainstaluj najnowszą wersję programu Thonny.
6. Uruchom program Thonny.
7. Kliknij na ikonkę trzech linii w prawym dolnym rogu i wybierz „Configure interpreter”.
8. Ustaw typ interpretera na „MicroPython (Raspberry Pi Pico)”.
9. Z menu w prawym dolnym rogu wybierz *MicroPython (Raspberry Pi Pico) Board CD @COMxx*. Interpreter w polu Shell z wyświetli informację o wersji: *MicroPython feature/psram-and-wifi, Pico2 v0.0.12 on 2025-02-28; Raspberry Pi Pico2 with RP2350*
10. Pobierz folder *code* z kodem aplikacji z repozytorium https://ep.com.pl/files/evo/13752-internet_rzeczy_w_pomiarach_srodowiskowych_21_pomiar_ph.zip
11. Otwórz w oknie *Files* folder *code*.
12. Kliknij prawym klawiszem myszy na plik *main.py* i wybierz *Upload to*.
13. Tak samo załaduj do płytki Pico 2 pozostałe pliki z folderu.

Oprogramowanie

Oprogramowanie przygotowane w języku Python jest wzorowane na bibliotece dostarczanej przez firmę DFRobot i przeznaczonej dla Arduino i języka Python [3]. Zostały przygotowane trzy pliki aplikacji.

- *Reset.py* – ustawianie domyślnych wartości parametrów przechowywanych w plikach *neutralVoltage.txt* oraz *acidVoltage.txt*. Początkowo (oraz po resetcie) są one ustawione na następujące wartości (mV): *neutralVoltage*=1500,00 (pH 7,0), *acidVoltage*=2032,44 (pH 4,0).
- *main.py* – wykonywanie standardowych pomiarów pH i temperatury. Temperatura jest mierzona za pomocą dwóch czujników DS18B20. Pomiar napięcia jest dokonywany przez przetwornik ADS1115 o zakresie napięć wejściowych 4096 mV, z ustawionym najdłuższym czasem pomiaru (częstotliwość próbkowania: 8 sps). Pomiar napięcia z sondy pH jest wykonywany dziesięciokrotnie. Następnie dane w buforze są sortowane, a 6 środkowych wyników jest uśrednianych. Na podstawie bieżącego pomiaru napięcia obliczane jest pH. Do obliczeń pobierane są z plików dane dwóch punktów kalibracyjnych. Rezultaty pomiarów trafiają do terminala oraz na wyświetlacz e-Paper.

Plik *Calibration* wykonuje kalibrację metodą dwupunktową i automatycznie identyfikuje dwa standardowe roztwory buforowe (4,0 i 7,0). Pomiary są wykonywane tak samo, jak w przypadku skryptu *main.py*. Wartości parametrów punktów kalibracji są zapisywane do plików *neutralVoltage.txt* oraz *acidVoltage.txt*. Opis postępowania w trakcie kalibracji jest zamieszczony w dalszej części opisu.

Konfiguracja pomiarowa

Konfiguracja pomiarowa (**fotografia tytułowa**) została skompletowana z użyciem płytki RPi Pico 2 [10], modułu DFRobot I²C ADS1115 [8] (zmodyfikowanego), dwóch sond temperatury z czujnikiem DS18B20, zestawu *Gravity: Lab Grade Analog pH Sensor Kit for Arduino/Raspberry Pi* (SEN0161-V2) [1] (zawierającego: sondę pH, płytkę pomiarową Gravity pH Meter V2.0, oraz buteleczki roztworów wzorcowych dla pH 4,0 i 7,0), wyświetlacza Pico Inky Pack (PIM634) oraz ekspandera Pico Omnibus (PIM556).

- Do płytki pomiarowej V2 dołącz kablem z wtyczką BNC sondę pomiarową.
- Do Pinu A0+ modułu DFRobot I²C ADS1115 dołącz (kabelkiem) pin (+) płytki pomiarowej V2, do pinu A0 dołącz pin (A), zaś do pinu A0- dołącz pin (-).
- Do płytki RPi Pico 2 podłącz moduł DFRobot I²C ADS1115: VDD do szyny VSYS, SCL do GP5, SDA do GP4 i GND do GND.
- Do płytki RPi Pico 2 dołącz sondy DS18B20: VCC (czerwony) do szyny 3,3 V, DATA (zielony) do GP16, GND (żółty) do GND.
- Płytkę RPi Pico 2 dołącz kablem USB do komputera.

Do obsługi oprogramowania w języku Python najłatwiej jest zastosować środowisko Thonny.

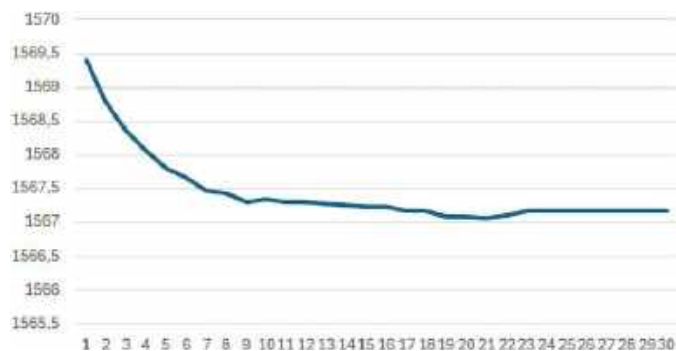
Kalibracja dwupunktowa czujnika pH

Kalibrację wielopunktową przeprowadza się z użyciem maksymalnie pięciu buforów standardowych. Użycie więcej niż pięciu punktów nie przynosi znaczącej poprawy uzyskiwanych informacji statystycznych. Funkcja kalibracji jest dana przez regresję liniową zmierzonych różnic potencjałów [15].

Kalibracja jednopunktowa jest niewystarczająca do określenia obu przesunięć nachylenia – charakteryzuje się największą niepewnością pomiaru pH, ponieważ zarówno nachylenie, jak i przesunięcie mogą zmieniać się wraz z wiekiem elektrod. Jest to zatem najmniej wiarygodna procedura.

W większości praktycznych zastosowań szklane ogniwa elektrodowe kalibruje się metodą dwupunktową, do której wymagane są dwa standardowe roztwory o stężeniach 4,0 i 7,0 (załączone w zestawie fabrycznym).

Ponieważ elektrody szklane są głównie przeznaczone do pomiaru pH w roztworach wodnych, w przypadku pomiaru rozpuszczalnika



Rysunek 3. Napięcie sondy przy pH 7,0



Rysunek 4. Napięcie sondy przy pH 4,0

innego niż woda odczyt pH będzie odbiegał od wartości oczekiwanych. W takich sytuacjach należy zbadać pomiar pH roztworu rozpuszczalnika mieszanego w laboratorium i porównać go z wynikiem uzyskanym w szklanej kuwecie przed rozpoczęciem właściwych pomiarów. Dodatkowo przed przystąpieniem do pracy w terenie należy wziąć pod uwagę stopień zużycia pierścienia uszczelniającego elektrody.

Aby zapewnić możliwie wysoką dokładność pomiarów, sondę pH należy skalibrować przed pierwszym użyciem oraz po dłuższym okresie nieużywania (najlepiej raz w miesiącu).

Pierwszy punkt kalibracyjny

- Umyj sondę wodą destylowaną, a następnie usuń resztki wody papierem.
- Umieść sondę pH w standardowym roztworze buforowym o stężeniu 7,0 i delikatnie zamieszaj.
- Uruchom program *Calibration.py*.
- Zatrzymaj pomiar po ustabilizowaniu się odczytów.

Drugi punkt kalibracyjny

- Umyj sondę wodą destylowaną, a następnie usuń resztki wody papierem.
- Umieść sondę pH w standardowym roztworze buforowym o stężeniu 4,0 i delikatnie zamieszaj.
- Uruchom skrypt *Calibration.py*.
- Zatrzymaj pomiar po ustabilizowaniu się wartości pomiarowych.

Po wykonaniu powyższych kroków kalibracja dwupunktowa jest zakończona, a następnie czujnik może zostać użyty do rzeczywistych pomiarów. Odpowiednie parametry procesu kalibracji zostały zapisane w plikach *neutralVoltage.txt* oraz *acidVoltage.txt*.

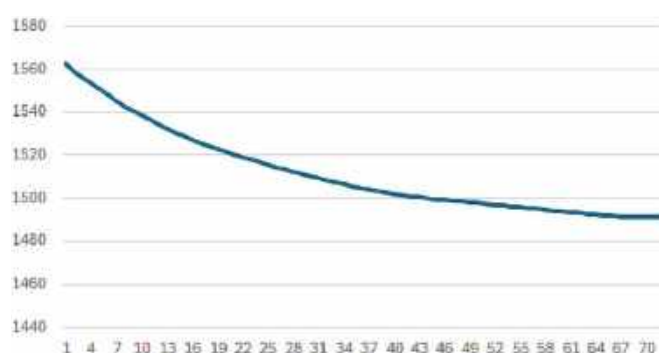
Pomiary pH

Zestaw uruchomieniowy *SHT31 Smart Gadget* firmy Sensirion to referencyjna płytka drukowana, która demonstruje dokładność, wydajność i łatwość użycia czujników wilgotności oraz temperatury SHT31 [13]. W ramach opisywanych eksperymentów zastosowano ją jako termometr odniesienia – pomiary temperatury dokonywane za pomocą obu czujników DS18B20 (położonych tak, żeby się stykały) oraz czujnika temperatury powietrza SHT31 wykazywały (po długim okresie stabilizacji) te same wartości, z dokładnością do drugiego miejsca po przecinku.

Wiarygodność danych pomiarowych zależy od dokładności i stabilności sondy pH. Dwa kluczowe czynniki to okresy stabilności po zmianie temperatury oraz wartości pH roztworu buforowego.

Przykładowe badanie sondy wyglądało następująco:

1. Elektrode przepłukano wodą dejonizowaną i umieszczono w buforze o pH 7,0. Sondę pozostawiono do ustabilizowania przez czas ok. 5 minut.
2. Elektrode przepłukano wodą dejonizowaną i przeniesiono do alikwotu buforu o pH 4,0 na okres ok. 5 minut.
3. Elektrode przepłukano wodą dejonizowaną i umieszczono w wodzie z kranu na okres ok. 10 minut.



Rysunek 5. Napięcie sondy przy pomiarze wody kranowej

Badania były przeprowadzone przy temperaturze roztworów ok. 23,44°C. Napięcie pomiarowe dla roztworu pH 7,0 dosyć szybko się ustabilizowało i pierwszy punkt pomiarowy został ustawiony na 1567,17 mV (**rysunek 3**). Pomiar roztworu o pH 4,0 wolniej się stabilizował i drugi punkt pomiarowy został arbitralnie ustawiony (po 10 minutach) na 2074,94 mV (**rysunek 4**). Producent podaje czas odpowiedzi poniżej 2 minut, co zgadza się tylko dla pierwszego punktu. Pomiar pH wody kranowej wykazywał wzrost pH (i wartości pomiarowej napięcia) od 7,03 na początku do 7,45 po 10 minutach oraz do 7,62 po 45 minutach (**rysunek 5**).

Typowy zakres pH w systemach wód powierzchniowych wynosi od 6,5 do 8,5, a w przypadku wód gruntowych – od 6,0 do 8,5.

Gwoli ścisłości należy dodać, że nakrętka butelki sondy została fabrycznie wyposażona w płyn ochronny (3,3 mol/l KCL), w którym czujnik powinien być przechowywany pomiędzy pomiarami. Jednak przy wykonywaniu pierwszych testów płyn ochronny zastosowanej w eksperymentach sondy został wylany i sensor przez ponad 6 miesięcy był przechowywany bez niego. Może to być przyczyną obserwowanego nadmiernego dryftu pomiarowego.

Pomiary zasilania

Do dynamicznego pomiaru prądu zasilania bardzo dobrze nadaje się zestaw Power Profiler Kit II (PPK2) firmy Nordic Semiconductor, opisany dokładniej w artykule „Profilowanie mocy z zastosowaniem Power Profiler Kit II” [18].

Pomiar był wykonywany przy zasilaniu z PPK2 szyny VBUS płytki RPi Pico 2 napięciem 5 V. Pobór mocy dla drugiego bloku pomiarowego, zarejestrowany po uruchomieniu programu, pokazano na **rysunku 6**. Kanały cyfrowe sygnalizują czas wykonywania operacji (poziom wysoki): CH0 – pomiar temperatury (589 ms), CH1 – pomiar napięcia z sondy pH (1,41 s), CH2 – przygotowanie (506 ms) i wyświetlanie danych (979 ms).

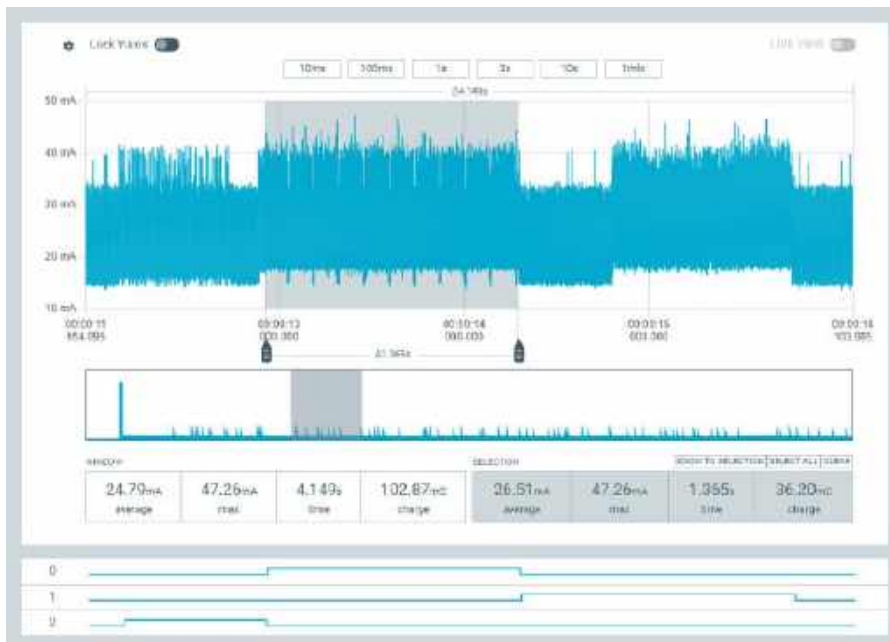
Przy włączaniu zasilania płytki Pico 2 pojawia się 0,88-milisekundowy pik prądu rozruchowego o maksymalnej wartości 0,67 A. Może to stwarzać problemy w przypadku źródeł zasilania małej mocy (np. odnawialnych).

Średni pobór prądu układu za okres 1 minuty wynosi 23,34 mA, a w czasie trwania bloku pomiarowego (3,6 s) 25,17 mA. Podczas pracy układu maksymalny pobór prądu nie przekracza 50 mA. Średni prąd pobierany przez procesor ma wartość rzędu 16,72 mA, a podczas operacji sleep – około 22,5 mA. Niestety implementacja MicroPythona na procesor RP2350 obecnie nie dostarcza obsługi trybów obniżonego poboru mocy, choć ta jest już dostępna w bibliotekach napisanych w języku C.

Podsumowanie

Czujnik pH DFRobot może być stosowany w różnych aplikacjach, takich jak akwapionika, akwakultura i badania środowiskowe wody. Zaprezentowany system pomiaru pH i temperatury pozwala na szybkie spełnienie szeregu wymagań funkcjonalnych.

W próbach został zastosowany jeden rdzeń Arm Cortex-M33 mikrokontrolera RP2350 oraz język MicroPython. Jednak usypianie rdzenia pomiędzy pomiarami okazało się niemożliwe z poziomu języka MicroPython. Pracujemy nad wdrożeniem użycia (w kodzie języka MicroPython) biblioteki języka C, oferującej taką funkcjonalność.



Rysunek 6. Prąd zasilania układu pomiarowego z sondą pH i wzmacniaczem

Warto dodać, że istnieje możliwość zastosowania płytki Raspberry Pico 2W zamiast Pico 2. Wtedy można obsługiwać zdalnie system z transmisją bezprzewodową, z użyciem protokołu Bluetooth lub Wi-Fi.

Opisane w artykule eksperymenty stanowią pierwsze podejście do zagadnienia pomiarów pH, które wymagają jednak przeprowadzenia szerszych badań. Dla poprawy jakości pomiarowej można też spróbować ekranowania płytki wzmacniacza pomiarowego oraz kabla doprowadzającego sygnał analogowy do przetwornika ADC.

Henryk A. Kowalski
Instytut Informatyki
Politechnika Warszawska

Literatura

- [1] Gravity: Lab Grade Analog pH Sensor Kit for Arduino/Raspberry Pi (with Calibration Solutions), DFRobot, <https://www.dfrobot.com/product-1782.html>
- [2] Gravity_Analog_pH_Sensor_Meter_Kit_V2_SKU_SEN0161-V2-DFRobot, DFRobot, https://wiki.dfrobot.com/Gravity_Analog_pH_Sensor_Meter_Kit_V2_SKU_SEN0161-V2
- [3] DFRobot_PH, https://github.com/DFRobot/DFRobot_PH
- [4] pH sensor Arduino, how do pH sensors work, application of pH meter, pH sensor calibration, Shahzada Fahad, Electroniclinic, <https://www.electroniclinic.com/ph-sensor-arduino-how-do-ph-sensors-work-application-of-ph-meter-ph-sensor-calibration/>
- [5] Skala pH, Wikipedia, https://pl.wikipedia.org/wiki/Skala_pH
- [6] Gravity: 7/24 Industrial Analog pH Meter Kit (Arduino, Raspberry Pi), <https://www.dfrobot.com/product-2069.html>
- [7] Waterproof DS18B20 Digital Temperature Sensor (DFR0198), DFRobot, <https://www.dfrobot.com/product-689.html>
- [8] Gravity: I²C ADS1115 16-Bit ADC Module, DFR0553, DFRobot, <https://www.dfrobot.com/product-1730.html>
- [9] ADS1115_mpy, A MicroPython module for the ADS1115 ADC. Wolfgang (Wolle) Ewald, https://github.com/wollewald/ADS1115_mpy
- [10] Raspberry Pi Pico 2, <https://www.raspberrypi.com/products/raspberry-pi-pico-2/>
- [11] Płytki Raspberry Pi Pico 2/2W z procesorem RP2350, Henryk A. Kowalski, EP 3/2025, <https://ep.com.pl/projekty/moduly-w-aplikacjach/16453-internet-rzeczy-w-pomiarach-srodowiskowych-15-plytki-raspberry-pi-pico-2-2w-z-procesorem-rp2350>
- [12] RP2350 Datasheet, 2024-10-16, Raspberry Pi, <https://datasheets.raspberrypi.com/rp2350/rp2350-datasheet.pdf>
- [13] SHT31 Smart Gadget Development Kit, Sensirion, <https://sensirion.com/products/catalog/SHT31-Smart-Gadget>
- [14] pH Sensor Reference Design Enabled for RF Wireless Transmission, Erbe D. Reyta, Mark O. Ochoco, Nov 1 2016, Analog Devices, <https://www.analog.com/en/resources/technical-articles/ph-sensor-reference-design-enabled-for-rf-wireless-transmission.html>
- [15] TI Designs – Wireless pH Sensor Transmitter (DevPack for SensorTag), June 2016, Texas Instruments, <https://www.ti.com/lit/ug/tidua47b/tidua47b.pdf>
- [16] Pico Graphics, Pimoroni, <https://github.com/pimoroni/pimoroni-pico/tree/main/micropython/modules/picographics>
- [17] Pimoroni Pico MicroPython for RP2350/Pico2 boards, <https://github.com/pimoroni/pimoroni-pico-rp2350>
- [18] Profilowanie mocy z zastosowaniem Power Profiler Kit II, EP 5/2022, <https://ep.com.pl/kursy/15267-systemy-dla-internetu-rzeczy-60-profilowanie-mocy-z-zastosowaniem-power-profiler-kit-ii>

AT-AD269S
Mikroskop cyfrowy
z ekranem 10 cali,
powiększenie do 5000×,
5 obiektywów i endoskop
ANDONSTAR AD269S-M



AT-AD409PRO
Mikroskop do lutowania
z profesjonalnym
metalowym stojakiem,
ekran 10,1 cala,
powiększenie do 300×, HDMI
ANDONSTAR AD409Pro



BESTSELLERY sklepu AVT – sklep.avt.pl

Mikroskopy cyfrowe dla elektroników

Rabat dla Czytelników EP
przy zakupie podaj kod **EP2505MC**

-3%

Rabat dla Prenumeratorów EP
przy zakupie podaj numer prenumeraty

-6%

AT-AD246S-M
Mikroskop cyfrowy 7 cali
z powiększeniem:
60...240×, 18...720×,
1560...2040×
ANDONSTAR AD246S-M



AT-AD407
Mikroskop cyfrowy 7 cali,
powiększenie do 270×
ANDONSTAR AD407



AT-AD249S-M
Mikroskop cyfrowy 10 cali
z powiększeniem:
60...240×, 18...720×, 1560...2040×
ANDONSTAR AD249S-M



AT-AD210
Mikroskop cyfrowy 5...260×
z wyświetlaczem 10,1 cala
ANDONSTAR AD210



QuantAsylum z serii QAxXX

– zaawansowany system pomiarowy audio

W cyklu artykułów dotyczących pomiarów amatorskich wzmacniaczy audio, opublikowanym na łamach EP 07...11/2024, w przystępny sposób zaprezentowano metody pomiaru elementów toru audio z użyciem wysokiej jakości karty dźwiękowej (interfejsu audio). Jest to najtańszy i najbardziej dostępny dla amatorów DIY sposób na sprawdzenie jakości konstruowanego lub serwisowanego sprzętu. W niniejszym artykule opiszemy natomiast zintegrowany system pomiarowy QuantAsylum – rozwiązanie opracowane z myślą o szeroko pojętych pomiarach torów audio.

Sprzętowy interfejs audio – wraz z oprogramowaniem Room Acoustics Software (REW) lub RightMark Audio Analyzer (RMAA) – daje spore możliwości oceny sprzętu. Sam korzystałem kiedyś z tego sposobu podczas uruchamiania, strojenia i testowania studyjnego sprzętu audio – używałem interfejsu MOTU828 i oprogramowania Adobe Audition, które – oprócz pomiarów – służyło też oczywiście do rejestracji i obróbki dźwięku. Nie była to najtańsza możliwa metoda, ale przy produkcji niskoseryjnej całkowicie akceptowalna kosztowo, w przeciwieństwie do rozwiązań w pełni profesjonalnych, takich jak np. analizator Audio Precision APx555 (**fotografia 1**). Sprzęt ten – w podstawowej wersji – dostępny jest w cenie przekraczającej 40 000 dolarów (!). Jest to w zasadzie jedno z nielicznych dostępnych rozwiązań do specjalistycznych pomiarów audio. Z kronikarskiego obowiązku i czystej sympatii dla systemu Hameg HM8000 warto przypomnieć moduły generatora o niskich zniekształceniach HM8037 i współpracującego z nim miernika zniekształceń nieliniowych HM8027, pokazane na **fotografiach 2a i 2b**, których używam od końca lat 90. ubiegłego wieku do dziś...

Oczywiście współczesne oscyloskopy o 12-bitowej rozdzielczości, z niskim poziomem szumów oraz dopracowaną analizą FFT (pomocną do orientacyjnej oceny poziomu zniekształceń) także nadają się

do pomiarów wstępnych – tym bardziej, że większość współczesnych modeli wspiera (wraz z współpracującym generatorem arbitralnym) automatyczne wykreślanie charakterystyk amplitudowo-fazowych Bodego, co daje nam możliwość szybkiej oceny pasma przenoszenia. W identyczny sposób do takich pomiarów można wykorzystać zestaw AnalogDiscovery3 z dostępnym pakietem Audio Analyzer Suite, ale pomimo większej rozdzielczości nie osiągniemy takiej jakości pomiarów, jak w przypadku karty dźwiękowej z przetwornikami 24-bitowymi, a nawet 32-bitowymi.

Istotnym ograniczeniem pomiarów pasma przenoszenia przy pomocy interfejsu audio jest maksymalna częstotliwość próbkowania równa 192 kHz – znacznie niższa, niż w przypadku oscyloskopu. Co prawda w nielicznych przypadkach karty dysponują próbkowaniem 352/384/768 kHz, ale tu już pojawiają się problemy z interfejsem komunikacyjnym (można tu w zasadzie zastosować tylko karty wewnętrzne PCIe) i stabilnym działaniem driverów programowych, szczególnie w przypadku produktów firm o niesprawdzonej renomie. A cena interfejsu renomowanej marki do najniższych niestety nie należy.

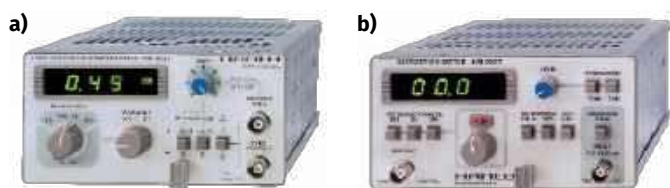
Istotnym problemem z punktu widzenia powtarzalności pomiarów jest płynna regulacja głośności, stosowana w niektórych interfejsach audio. Z takim rozwiązaniem wiąże się brak kalibracji poziomów sygnałów i wbudowanych tłumików – które dopasowane są do maksymalnych poziomów typowych w torach audio. Brak wbudowanych, skalibrowanych dzielników pomiarowych i stopni wyjściowych o większej wydajności prądowej także zawęża zakres możliwych do wykonania pomiarów. Jeżeli realizujemy pojedyncze pomiary, możemy pogodzić się z tymi niedogodnościami, oszczędzając nieco gotówki – a w zasadzie nie wydając jej wcale, bo kartę dźwiękową lub interfejs audio ma raczej każdy elektronik zainteresowany techniką audio i z pewnością nie mówimy tutaj o podłej jakości karcie USB. A i oprogramowanie udostępniane jest nieodpłatnie.

Jeżeli jednak wykonujemy większą liczbę pomiarów (w tym także powtarzalnych), a nie mamy zamiaru skupiać się na każdorazowym przygotowaniu i kalibrowaniu stanowiska oraz opracowaniu kilku brakujących elementów toru pomiarowego (takich jak dzielniki, wzmacniacze pomiarowe i moduły obciążenia), ani lub tym bardziej wydawać znacznych środków na sprzęt Audio Precision, warto przyglądnąć się rozwiązaniom amerykańskiej firmy QuantAsylum. Marka ta – borykając się z podobnymi problemami – opracowała modułowy zestaw pomiarowy przeznaczony stricte do pomiarów audio.

Zestaw pomiarowy składa się z kilku niezależnych modułów, o które można uzupełniać warsztat sukcesywnie, w zależności od typu wykonywanych pomiarów. Sercem i niezbędnym elementem zestawu pomiarowego jest analizator audio QA403 pokazany na **fotografii 3**, wraz z darmowym oprogramowaniem QA40x



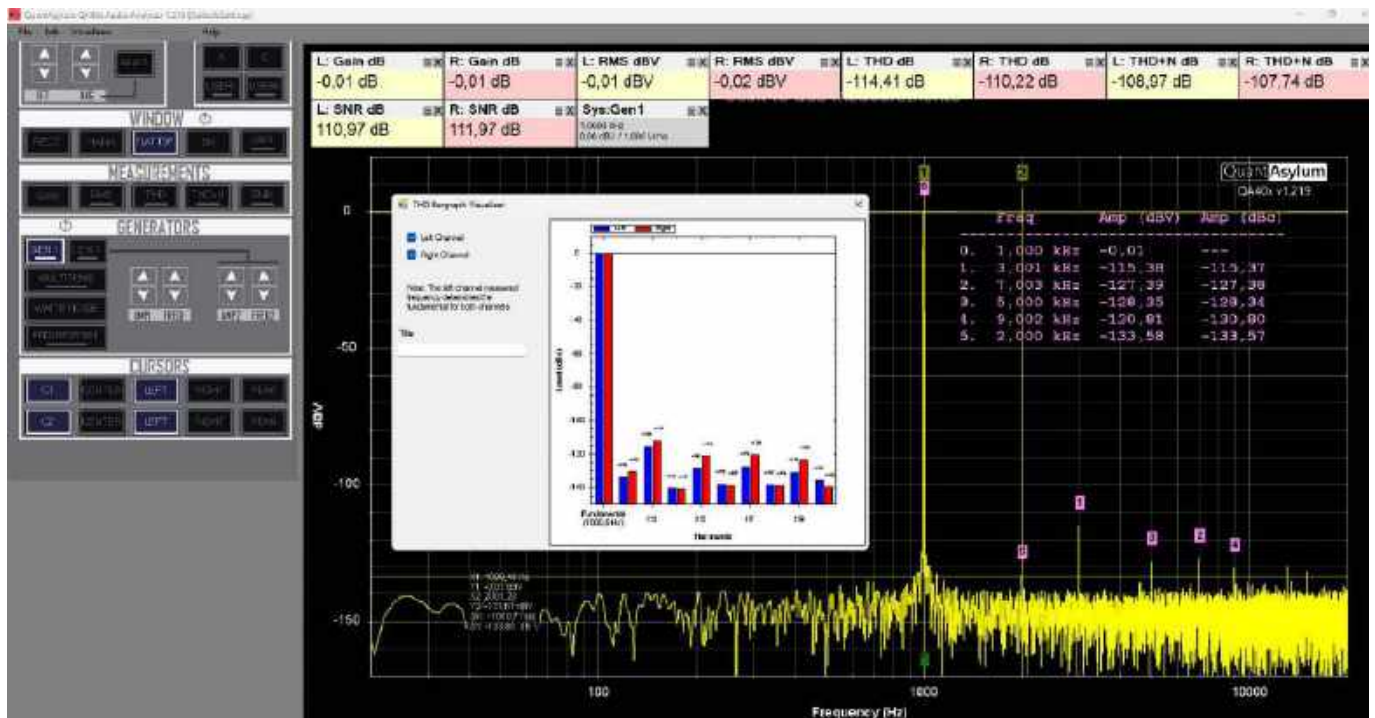
Fotografia 1. Analizator audio APx555 (fot. z materiałów Audio Precision)



Fotografia 2a,b. Nieprodukowane już moduły generatora o niskich zniekształceniach HM8037 i miernika zniekształceń nieliniowych HM8027 firmy Hameg (dzisiaj Rohde & Schwarz)



Fotografia 3. Analizator audio QA403 (z materiałów QuantAsylum)



Rysunek 1. Oprogramowanie sterujące analizatorem QA403

niewymagającym kluczy sprzętowych lub rejestracji oraz pozbawionym restrykcji pod względem ilości zainstalowanych instancji. Software pracuje na systemach Windows 10/11 oraz Linux.

Jest to już czwarta wersja analizatora o całkiem sporych możliwościach pomiarowych. Jak zapewniają projektanci, nie mają oni zamiaru rywalizować z analizatorami przekraczającymi wartość samochodu, ale w 99% przypadków nie ma nawet takiej potrzeby. Dodatkowo niewielkie rozmiary (177×97×44 mm) i elastyczność oprogramowania znacznie ułatwiają testowanie urządzeń audio, nie tylko w warsztacie, ale także w warunkach polowych – a to za sprawą zasilania i komunikacji poprzez izolowany port USB (maksymalny pobór prądu to 900 mA). Analizator jest tak naprawdę wysokiej jakości, dwukanałowym interfejsem audio o rozdzielczości 32 bitów, maksymalnej częstotliwości próbkowania 192 kHz i uzupełnionym o wbudowane, kalibrowane wzmacniacze wejściowe i wyjściowe z tłumikami. Analizator pozbawiony jest jakichkolwiek elementów manipulacyjnych, gdyż obsługa odbywa się tylko programowo. Na froncie urządzenia umieszczono gniazda BNC, na które wyprowadzono dwa kanały symetrycznego interfejsu wejściowego i wyjściowego oraz diody sygnalizujące zasilanie, połączenie z oprogramowaniem i aktywny tłumik. Gniazda BNC w technice audio nie są stosowane zbyt często i wielu użytkowników wolałoby typowe gniazda XLR, jack 6,3 mm lub RCA, ale nie jest to problem spędzający sen z powiek. Impedancja wejściowa w trybie symetrycznym wynosi 200 kΩ, w trybie niesymetrycznym 100 kΩ. Podczas pomiarów w trybie niesymetrycznym, dla zachowania niskiego poziomu zakłóceń, należy zewrzeć wejścia „–”. Można zastosować w tym celu terminator BNC50W, nieznacznie pogarszając stosunek sygnał-szum (osobiście używam przerobionych terminatorów 50 Ω, w którym wbudowany rezystor 50 Ω zastąpiłem zworą). Maksymalne napięcie (AC+DC) na wejściu analizatora nie powinno przekraczać 56 Vpp, użyteczny zakres pomiarowy to +18 dBV (8 Vrms) przy wyłączonym tłumiku oraz +42 dBV przy tłumiku włączonym – jest on jednak ograniczony programowo do +32 dBV (40 Vrms). Wejście sprzężone jest zmiennoprądowo i przenosi sygnały od 1 Hz w górę. Poziom

szumów wynosi –115 dBV przy zwartym wejściu i $f_s=48$ kHz, poziom zniekształceń THD+N jest równy –105 dBV. Sygnał wyjściowy jest sprzężony stałoprądowo, impedancja wyjścia wynosi odpowiednio 200 Ω w trybie symetrycznym i 100 Ω w trybie niesymetrycznym (zalecane obciążenia >1 kΩ), a maksymalny poziom sygnału wyjściowego to +18 dBV (8 Vrms).

QA403 ma wbudowane automatyczne tłumiki 0, 10, 20 i 30 dB. Poziom szumów wynosi –105 dB przy sygnale wyjściowym 0 dBV (–100 dBV/15 dBV), a zniekształcenia THD+N utrzymują się na poziomie –105 dB. Dodatkowe złącze Expansion Connector umożliwia wyprowadzenie izolowanego interfejsu I²S, np. do testów przetworników D/A. Na tylnej ścianie znajduje się tylko port komunikacyjny i zasilający USB z gniazdem typu A (przydałoby się już złącze typu C).

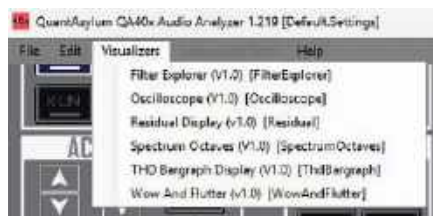
Pod względem konstrukcyjnym analizator – jak zapewnia producent – nie jest zlepką przypadkowych modułów, lecz opracowanym od podstaw systemem pomiarowym. Dobrze o tym wiedzieć, patrząc na jego aktualną cenę wynoszącą 600 dolarów... Całość zmontowana jest na jednej czterowarstwowej płytce drukowanej,

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

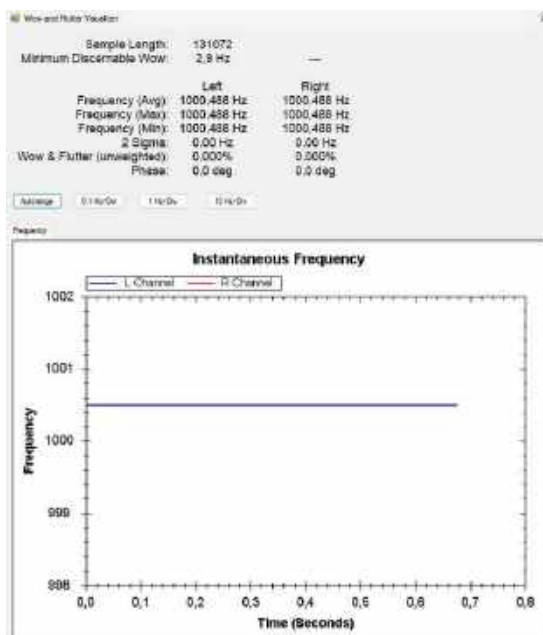


Rysunek 2. Narzędzia do analizy graficznej

głównym kontrolerem odpowiadającym za komunikację USB oraz obsługę wbudowanych przetworników A/D i D/A jest mikrokontroler LPC55S16 firmy NXP (ARM Cortex-M33). Tor wejściowy korzysta z 32-bitowego przetwornika A/D Sabre typu ESS9822 z buforami opartymi na OPA1612 i wzmacniaczu różnicowym OPA1632. Tor wyjściowy, z 32-bitowym przetwornikiem D/A Sabre typu ESS9038, służy natomiast jako generator przebiegów testowych. Jak widać, nie ma tu niepotrzebnych oszczędności, a zastosowane układy należą do czołówki współczesnych rozwiązań audio. Wsparcie użytkownika udzielane jest na forum: <https://forum.quantasylum.com/> a samo oprogramowanie można pobrać ze strony: <https://quantasylum.com/pages/downloads-1>

Oprogramowanie ma nieco nietypowy wygląd, ale nie przeszkadza to w użytkowaniu, chociaż przydałoby się trochę usprawnień, dotyczących m.in. szybkiego wyboru wartości amplitudy i częstotliwości generatorów. Przykładowy zrzut ekranu podczas pomiaru pętli A/D/D/A QA403 pokazano na **rysunku 1**. Na ekranie mogą być wyświetlane aktualne pomiary parametrów, takich jak poziomy sygnał, zniekształcenia czy stosunek sygnał-szum, konfigurowane prawym przyciskiem myszy w obszarze pomiarów. Możliwy jest pomiar z użyciem kursorów, po zaznaczeniu szczytu sygnału wyświetlona zostanie tabela z bieżącymi wartościami. Wbudowane narzędzia wybierane z menu Visualisers, takie jak Filter Explorer, Oscilloscope, Residual Display, Spectrum Octaves, THD Bargraph Display czy Wow and Flutter, pokazane na **rysunku 2**, umożliwiają graficzną ocenę parametrów oraz eksport wykresów do plików graficznych w typowych formatach.

Szczególnie ostatnie narzędzie (**rysunek 3**), dziś już nieczęsto używane, może okazać się przydatne dla konstruktorów gramofonów DIY oraz osób przywracających do użytku sprzęt vintage, taki jak gramofony, magnetofony szpulowe i kasetowe, w ocenie jakości pracy napędu



Rysunek 3. Ocena nierównomierności przesuwu

(pomiar prędkości i nierównomierności odtwarzania), przy użyciu odpowiednio płyty lub taśmy testowej, analogicznie jak dawniej przy użyciu miernika nierównomierności napędu (np. Hameg HM8026). W przykładzie z rysunku 3 pokazano wynik pomiaru stabilności wewnętrznego generatora, jakiej należało się spodziewać – sygnał odtworzony z płyty lub taśmy nie będzie już niestety idealną prostą.

Ekran aplikacji może zostać skonfigurowany w domenie czasu (działając jak oscyloskop) lub domenie częstotliwości (działając jak analizator widma). Udostępniona jest oczywiście konfiguracja częstotliwości próbkowania, parametrów FFT i uśredniania, a także możliwość wyboru krzywych pomiarowych ważonych typu A i C oraz krzywych definiowanych przez użytkownika, a przydatnych przy pomiarach np. korektorów RIAA. Tor wyjściowy z przetwornikami D/A służy do generowania typowych sygnałów używanych podczas pomiarów, w tym:

- sinusoidy (dwa sumowane kanały z niezależnie regulowanymi częstotliwościami i amplitudami, przydatne przy pomiarach intermodulacji),
- przebiegu sinusoidalnego multitone (wieloczęstotliwościowy z konfigurowaną ilością tonów w oktawie i całkowitym poziomie RMS),
- szumu białego,
- sygnału expo chirp (wykładnicze przemiatanie częstotliwości, z możliwą maską użytkownika do kwalifikacji sprawny/uszkodzony).

Konfiguracja podlegają też tłumiki sygnału wejściowego w zakresie 0...42 dB, ustawiane co 6 dB. Analizator pracuje w dwóch trybach: domyślnej akwizycji danych (która uruchamiana jest przyciskiem RUN) oraz IDLE, w którym dane nie są zbierane, a aktywny pozostaje tylko generator. W tym ostatnim trybie można zweryfikować wartości napięć w układzie przy pomocy multimetru lub oscyloskopu. Z menu RUN/STOP możemy także wybrać odtworzenie pliku *.wav (24/32-bit, float, z fs zgodną z ustawioną) do subiektywnej oceny jakości badanego urządzenia, a przy użyciu wyjścia I²S – do oceny konstruowanych przetworników D/A.

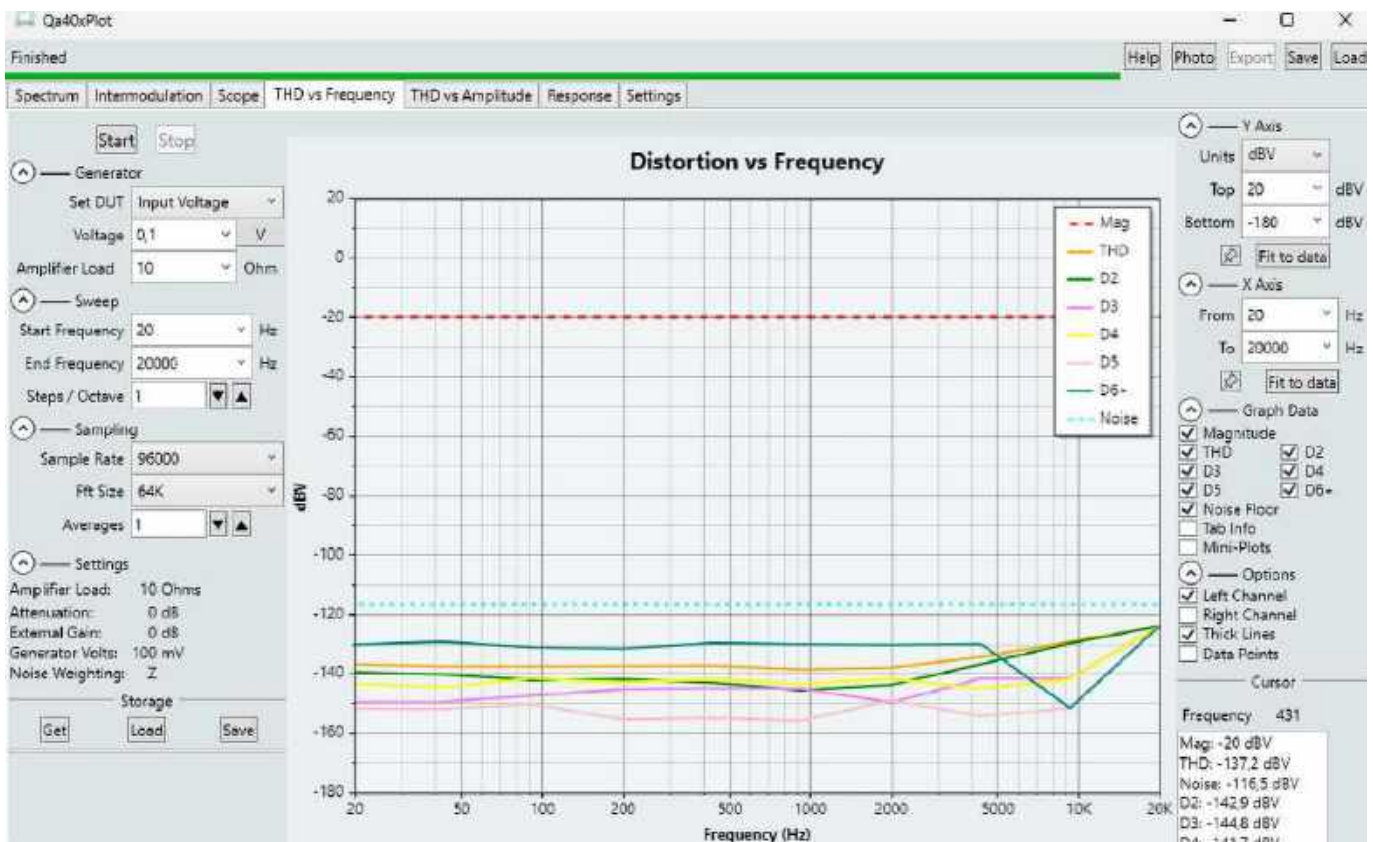
Parametry, ustawienia aplikacji oraz dane mogą zostać zapisane do plików konfiguracji lub danych i użyte w celu dalszej analizy lub szybkiej konfiguracji QA403 poleceniami export/import.

Oprócz udostępnianego przez producenta oprogramowania, analizator obsługuje także komendy HTTP GET/PUT oraz interfejs REST API, co umożliwia automatyzację pomiarów w różnych środowiskach, także w Matlabie, pobierając z analizatora „surowe” dane do dalszej analizy we własnym oprogramowaniu. Software jest aktualizowany, więc co jakiś czas warto sprawdzać, czy przypadkiem nie rozbudowano jego funkcjonalności. Warto poświęcić też uwagę instrukcji i zestawowi ustawień usprawniających typowe pomiary, przygotowanym przez Boba Cordella (autora doskonałych podręczników *Designing Audio Circuits and Systems* i *Designing Audio Power Amplifiers*). Materiały te znajdują się pod adresem:

<https://cordellaudio.com/instrumentation/quantasylum.shtml>



Fotografia 4. Programowalne obciążenie QA451 (fot. z materiałów QuantAsylum)



Rysunek 4. Alternatywne oprogramowanie Qa40xPlot

Ze względu na „otwartość” sterowania i akwizycji dostępne są też alternatywne wersje oprogramowania, jak Qa40xPlot (**rysunek 4**) dostępny pod adresem: <https://mzschmann.github.io/>

Drugim elementem zestawu niezbędnym podczas pomiaru wzmacniaczy mocy jest programowalne obciążenie QA451, zaprezentowane na **fotografii 4**.

W skład dwukanałowego obciążenia, oprócz sterowanych rezystorów 4/8 Ω , wchodzi tłumik 12 dB z filtrem dolnoprzepustowym 6. rzędu o $f_c=67$ kHz, umożliwiające pomiary końcówek mocy pracujących w klasie D. Maksymalny sygnał wejściowy to 40 Vrms, a maksymalny prąd ma wartość 10 A – co umożliwia pomiary mocy do ok. 200 W. Należy zwrócić uwagę, że pomiary odbywają się przy użyciu sygnału chirp z szybkim przemiataniem, aby nie przekroczyć temperatury rezystorów obciążenia.

Uwaga: QA451 nie nadaje się do długotrwałych testów wzmacniacza pod obciążeniem!

Dodatkowym wyposażeniem QA451 jest moduł pomiaru prądu pobieranego przez badany wzmacniacz, a także sterowany klucz zasilania o maksymalnych parametrach łączeniowych 50 VDC/10 A z funkcją soft-start. Klucz ten jest użyteczny przy pomiarze pobieranego przez wzmacniacz prądu, ale tylko przy zasilaniu niesymetrycznym – co niestety nie zawsze jest warunkiem spełnionym w przypadku badanych konstrukcji.

W zależności od wersji (QA451A lub QA451B) mamy do czynienia z różną rozdzielczością pomiaru prądu zasilania. Wersja A oferuje rozdzielczość 5...10 mA, a wersja B – 10...20 mA. W praktyce nie ma to większego znaczenia, chyba że potrzebujemy wyznaczyć sprawność z bardzo dużą dokładnością. Niestety chwilowo problem wyboru wersji jest arbitralnie rozwiązany, gdyż producent boryka się z problemami

z dostępnością elementów do QA451. Gdy moduł był osiągalny, w sklepie producenta kosztował 319 dolarów. Urządzenie jest zasilane i sterowane – podobnie jak QA403 – z izolowanego funkcjonalnie interfejsu USB-A (± 100 V), pobór prądu może sięgać 500 mA. Obciążenie elektroniczne ma identyczne wymiary, jak analizator, co umożliwia ich łączenie w systemie 1U, 2U przy pomocy dostępnych adapterów montażowych. Aktualnie dostępny jest uchwyt montażowy 1U w cenie 50 dolarów (!?). Podobnie jak QA403, obciążenie sterowane jest przez udostępnione oprogramowanie QA451 (**rysunek 5**) oraz interfejs REST.

Aplikacja umożliwia sterowanie zasilaniem badanego wzmacniacza i rezystorami obciążenia oraz mierzy ich temperaturę i prąd pobierany przez wzmacniacz. Szkoda, że nie zostały udostępnione gniazda do sterowania zewnętrznymi rezystorami mocy oraz że pomiar prądu dokonywany jest tylko w jednej gałęzi zasilania i pozbawiony funkcji pomiaru napięcia, aby można było bezpośrednio szacować pobieraną moc i sprawność badanego wzmacniacza. Problem pomiaru w przypadku większych mocy wzmacniaczy z zastosowaniem zewnętrznego rezystora obciążenia omówiony został na blogu dostępnym pod adresem: <https://quantasylum.com/blogs/news/using-the-qa451-with-external-loads>

Posiadając pierwsze dwa moduły, możemy przetestować i pomierzyć większość typowych rozwiązań wzmacniaczy i przedwzmacniaczy audio.



Rysunek 5. Aplikacja sterująca programowalnym obciążeniem QA451



Fotografia 5. Przedwzmacniacz niskoszumny



Fotografia 6. Driver szerokopasmowy

Jeżeli planujemy pomiary układów o niskich poziomach sygnału lub wymagających wysokiej impedancji obciążenia, niezbędnym elementem toru pomiarowego okaże się dwukanałowy, niskoszumny przedwzmacniacz QA472, zaprezentowany na **fotografii 5**. W obudowie o wymiarach identycznych jak QA403 i QA451 umieszczono dwa niezależne stopnie wzmacnienia o różnym przeznaczeniu.

Pierwszy z nich – PRE1 – przystosowany jest do podłączenia kalibrowanego mikrofonu pomiarowego, np. podczas pomiaru charakterystyk przenoszenia zestawów głośnikowych lub charakterystyki pomieszczenia. Wejście akceptuje symetryczny sygnał doprowadzony do gniazd BNC lub XLR. PRE1 oferuje ustalone skokowo (przyciskiem Gain) wzmacnienie 0/10/20 dB. Dodatkowo wejście XLR ma odłączany przyciskiem +48 V zasilacz Phantom. Impedancja wejściowa wynosi 100 k Ω w trybie zbalansowanym i 55 k Ω w niezbalansowanym. Przy wyłączonym lub włączonym zasilaniu Phantom wartości te wynoszą odpowiednio 13,6 k Ω i 6,8 k Ω . Do budowy PRE1 użyto precyzyjnego, niskoszumnego wzmacniacza różnicowego INA849. Stosunek sygnał-szum wynosi w najgorszym przypadku –100 dBV (przy największym wzmacnieniu 20 dB), zniekształcenia plasują się poniżej –110 dB/1 kHz, pasmo przenoszenia sięga 2,5 MHz, a impedancja wyjściowa wynosi 100 Ω . Drugi przedwzmacniacz PRE2 ma stałe wzmacnienie 30 dB i wysoką impedancję wejściową równą 1 M Ω . Szum jest niższy niż 102 dBV, zniekształcenia nie przekraczają 90 dBV, a pasmo przenoszenia sięga 250 kHz (ograniczone jest przez QA403, co nie dotyczy zastosowań niezależnych QA472). Przedwzmacniacz zasilany jest z izolowanego portu USB typu A, ale nie jest przez niego sterowany. Koszt urządzenia to 279 dolarów.

Ostatnim, czwartym elementem systemu jest driver przetworników QA462, którego wygląd pokazano na **fotografii 6**. QA472 to szerokopasmowy (500 kHz) wzmacniacz o niewielkiej mocy i wzmacnieniu 20 dB, zdolny dostarczyć do obciążenia prąd o natężeniu do 1,5 A. QA472 umożliwiaysterowanie obciążenia typu głośnik lub innego, o „trudnym” charakterze, z zapewnieniem stabilności. Driver wyposażono w obwód pomiaru prądu ze wzmacnieniem 50 $\times$ i konfigurowanym za pomocą przycisku rezystorem pomiarowym 1/10/100 Ω . Nie zabrakło także obwodu pomiaru napięcia



Fotografia 7. Kompletny system pomiarowy audio QuantAsylum

na podłączonym obciążeniu, wyposażonego we wbudowany tłumik –20 dB, umożliwiającym – opróczysterowania głośnika w celu pomiaru jego pasma przenoszenia – także pomiar impedancji złożonych układów zwrotnic głośnikowych lub elementów RLC. W rzeczywistości rezystory pomiarowe mają wartości 0,02/0,2/2 Ω , ale po 50-krotnym wzmacnieniu odpowiada to 1/10/100 Ω . Obwód niskoszumnego wzmacniacza 20 dB o impedancji wejściowej 50 k Ω wykonano z użyciem OPA1655, stopień drivera z konfigurowanym za pomocą przycisku ograniczeniem prądowym 300 mA/1500 mA zbudowano w oparciu o OPA564, a obwód pomiaru prądu korzysta ze wzmacniacza różnicowego INA849. Poziom szumów nie przekracza 20 μ V(A), a zniekształcenia są mniejsze niż 80 dB w warunkach: 1 kHz/1 W/8 Ω . Szczegółowe parametry i zbiory charakterystyk zamieszczono w karcie katalogowej.

Driver umieszczony jest w obudowie zgodnej z pozostałymi elementami. Ze względu na pobór mocy dostarczany jest z zewnętrznym zasilaczem 12 V/1,5 A. Koszt drivera wraz z zasilaczem to 329 dolarów.

Kompletny system pomiarowy – wraz z pomocniczymi akcesoriami i adapterami BNC/Mini XLR oraz I²S własnej konstrukcji – można zobaczyć na **fotografii 7**. Cena całości zawiera się w ok. 1530 dolarów (bez opłat dodatkowych i kosztów transportu). Jest to koszt zbliżony do ceny przyzwoitego oscyloskopu, a kwota wydaje się być akceptowalna dla mniejszych firm i manufaktur audio – zwłaszcza, że podział na moduły umożliwia zakup tylko niezbędnych elementów systemu lub jego stopniową komplektację.

Niestety dystrybucja w Europie jest bardzo słaba i najlepiej robić zakupy bezpośrednio u producenta, co wiąże się z pewnymi formalnościami i stosunkowo dużymi kosztami transportu. O ile jeszcze zdążymy go kupić przed wprowadzeniem kolejnych cel...

Adam Tatuś, EP

TAWOIA Glass (szkło kwarcowe)

<https://sklep.avt.pl/pl/menu/tawoia-glass-4505.html>



BESTSELLERY sklepu AVT – sklep.avt.pl

3 unikalne
serie
gniazdek
i
włączników

Rabat dla Czytelników EP
przy zakupie podaj kod **EP2505GW**

-5%

Rabat dla Prenumeratorów EP
przy zakupie podaj numer prenumeraty

-10%

Ceramic Loft (ceramika)

<https://sklep.avt.pl/pl/menu/seria-ceramic-loft-4190.html>



Retro PRL (bakelit)

<https://sklep.avt.pl/pl/series/retro-prl-3237.html>





Synteza dźwięku (1)

Informacje podstawowe

Od zarania dziejów aż do XX wieku cała działalność muzyczna ludzkości opierała się na wykorzystaniu naturalnych zjawisk i mechanizmów do wytwarzania dźwięku. Wraz z powstaniem elektroniki, a konkretniej pierwszych wzmacniaczy i oscylatorów, ludzkość zyskała nową metodę tworzenia brzmień – syntezę dźwięku.

nerowania dźwięków metodą naturalną są różnorodne. Najprostszy przykład to użycie możliwości ludzkiego narządu mowy do wytwarzania dźwięków o złożonej strukturze melodycznej. Z kolei uderzając jednym przedmiotem o inny można emitować nie tylko proste dźwięki perkusyjne, ale też złożone i bogate tony, wykorzystując różnorodne materiały: od wydrążonej kłody z naciągniętą na otwór membraną ze skóry zwierzęcej, aż po złożone, geometrycznie kształty z metalu czy nawet metalowe rury i płyty. W użyciu nadal pozostają też rury o różnych długościach, w których słup powietrza jest wprawiany w drgania, a długość i średnica rury określają częstotliwość rezonansową. Inną, wszechstronnie stosowaną metodą jest wprawianie w drgania napiętego włókna naturalnego lub wykonanej z metalu bądź tworzywa struny, a drgania te są następnie wzmacniane za pomocą pudła rezonansowego. Choć we wszystkich opisanych przypadkach ograniczamy się zaledwie do kilku bazowych metod, to bogactwo brzmienia jest ogromne i zależy od fizycznych charakterystyk konstrukcji instrumentów oraz użytych materiałów. Gitara klasyczna, banjo i mandolina to zasadniczo ten sam typ chordofonu szarpanego, a jednak każdy z tych instrumentów brzmi zupełnie inaczej. Podobnie jest w przypadku

syntezy dźwięku: mając do dyspozycji kilka podstawowych typów syntezy i „klocków” można wytworzyć najróżniejsze dźwięki. W ramach nowej serii artykułów przyjrzymy się budowie i zasadom działania syntezy dźwięku oraz ich elementów składowych. Pierwsza część cyklu została poświęcona omówieniu podstawowych zagadnień związanych z syntezą dźwięku.

Czym jest synteza dźwięku?

Synteza dźwięku jest procesem, w którym złożone brzmienie tworzone jest z prostszych sygnałów zmiennych, generowanych na drodze elektronicznej, a następnie poddanych procesowi obróbki, również za pomocą odpowiednich układów lub oprogramowania. Częstotliwość wytwarzanego dźwięku kontrolowana jest zazwyczaj za pomocą klawiatury sterującej lub sekwencera, czyli wcześniej zaprogramowanego automatu odtwarzającego sekwencję sygnałów sterujących syntezatorem (bądź



Fotografia 1. Syntezytor monolityczny Solina pozwalający tworzyć brzmienia naśladujące sekcję smyczkową orkiestry (<https://t.ly/yPIFs>)



Fotografia 2. Syntezator modularny Moog System 55, edycja limitowana z 2015 roku. Brzmienie jest tworzone przez łączenie kablami poszczególnych modułów i regulowanie ich nastaw. Im więcej dostępnych modułów w tego typu syntezatorze, tym bardziej rozbudowane i różnorodne brzmienia można tworzyć za jego pomocą (https://t.ly/sp_0G)

syntezatorami). Sam syntezator może też być kontrolowany innym instrumentem muzycznym – przez próbkowanie jego brzmienia albo przez protokół MIDI. Pod względem konstrukcji syntezatory można podzielić na trzy kategorie:

- syntezatory monolityczne,
- syntezatory modularne,
- syntezatory wirtualne.

Syntezatory monolityczne zawierają określone elementy niezbędne do wytwarzania dźwięku, połączone ze sobą w jeden, określony przez konstruktora sposób. Użytkownik ma kontrolę nad częścią lub całością parametrów, ale nie może zmienić kombinacji i układu połączeń elementów syntezatora. **Fotografia 1** prezentuje syntezator brzmień podobnych do sekcji skrzypcowej orkiestry – model Solina.

Syntezator modularny zawiera od kilku do kilkudziesięciu niezależnych od siebie elementów: oscylatorów, modulatorów, filtrów, generatorów obwiedni i wzmacniaczy sterowanych napięciem, które użytkownik łączy ze sobą w dowolny sposób za pomocą przewodów. Każdy moduł ma szereg regulatorów, które można kontrolować ręcznie lub za pomocą sygnałów z innych modułów. W ten sposób użytkownik ma pełnię kontroli nad tworzoną dźwiękiem, ale kosztem większej złożoności – nie tylko syntezatora, ale i samego procesu przygotowania nowego brzmienia. Na **fotografii 2** pokazano syntezator Moog System 55, produkowany w latach 1973...1981, a także – w limitowanej liczbie egzemplarzy – w 2015 roku.

Syntezatory wirtualne istnieją w formie oprogramowania i korzystają z technik DSP do modelowania brzmienia fizycznego syntezatora. Oprogramowanie takie może działać na komputerze PC w formie niezależnego programu, wtyczki VSTi (ang. Virtual Studio Technology instrument) do programu DAW (ang. Digital Audio Workstation) lub na dedykowanym sprzęcie, jako firmware mikrokontrolera/procesora aplikacyjnego lub układu FPGA. Syntezator wirtualny może być zarówno odwzorowaniem istniejącego syntezatora klasycznego, jak i zupełnie nowatorskim projektem. **Fotografia 3a** pokazuje przykładowy syntezator Yamaha Reface CS, będący współczesnym, wirtualnym syntezatorem



Fotografia 3. Syntezator Yamaha Reface CS (a) jest wirtualnym syntezatorem analogowym, mającym uchwycić brzmienie i charakter syntezatora Yamaha CS-80 (b). Nie jest to jednak stuprocentowa replika, a bardziej hołd złożony oryginałowi (<https://t.ly/DqhMY>, <https://t.ly/F4icy>)

analogowym, inspirowanym konstrukcją monolitycznego modelu Yamaha CS-80 (**fotografia 3b**).

Wszystkie syntezatory można też podzielić pod względem metody generowania dźwięku na:

- addytywne,
- subtraktywne,
- FM,
- wavetable,
- granularne,
- oparte na zniekształcaniu fazy,
- pozostałe (niesklasyfikowane).

Synteza addytywna opiera się na założeniu Fouriera, w myśl którego każdy złożony przebieg okresowy można stworzyć sumując przebiegi sinusoidalne o różnych częstotliwościach, fazach i amplitudach. Kilka oscylatorów kontrolowanych napięciem wytwarza przebiegi o określonych relacjach częstotliwościowych i fazowych, a ich sygnały są sumowane, tworząc złożony sygnał, który poddawany jest następnie dalszej obróbce.

W syntezie subtraktywnej złożony sygnał audio z oscylatora lub bloku oscylatorów (addytywnego) jest przepuszczany przez różne filtry sterowane napięciem i wzmacniacz (także przestrajany napięciowo), kontrolowane za pomocą sygnałów o niskiej częstotliwości lub/i sygnałów z generatora obwiedni. Od złożonego sygnału odejmowane są elementy niepożądane, a wzmacnieniu ulegają te, na których zależy użytkownikowi. Elementy syntezy subtraktywnej spotyka się w innych typach syntezatorów, gdyż jest to prosta metoda na uzyskanie złożonego brzmienia.

W przypadku syntezy FM dwa (lub więcej) oscylatory pracują synchronicznie z różnymi częstotliwościami. Sygnał jednego oscylatora moduluje częstotliwość drugiego, który może z kolei sterować częstotliwością kolejnego. Ten ostatni może być także sumowany z przebiegiem wyjściowym innego bloku syntezy FM. W efekcie powstaje złożony i bogaty harmonicznie dźwięk, który – podobnie jak w innych rodzajach syntezatorów – można poddać dalszej obróbce. Zmieniając stosunek częstotliwości fali nośnej do sygnału



Fotografia 4. Minisyntezator Stylophone z lat 70. ub. wieku w trakcie gry (<https://t.ly/O4fLb>)

modulującego da się uzyskać różne brzmienia, w tym gongi, dzwony rurowe czy inne dźwięki perkusyjne. Synteza FM była popularną metodą tworzenia dźwięków w komputerach i konsolach do gier w latach 80., a firma Yamaha wypuściła na rynek szereg dedykowanych układów scalonych (na przykład YM2151). Warto tutaj nadmienić, że opisana forma syntezy jest prosta do realizacji drogą cyfrową i z tego wynika jej popularność.

Synteza typu wavetable przypomina swoim działaniem układy DDS. W pamięci zapisana jest tablica dyskretnych wartości dla wybranych przebiegów. W trakcie syntezy kolejne wartości są odczytywane i wysyłane do przetwornika cyfrowo-analogowego. Częstotliwość taktowania tego procesu określa częstotliwość dźwięku wyjściowego. Dodatkowo wartości z tablicy można poddać innym procesom (np. modulacji FM) celem uzyskania jeszcze bardziej złożonego brzmienia. Syntezatory wektorowe używają dwóch (lub więcej) tablic z różnymi przebiegami, a ich wartości są ze sobą mieszane w stosunku wybranym przez użytkownika, najczęściej za pomocą joysticka.

Synteza granularna przypomina nieco metodę wavetable, z jedną różnicą: większa tablica, zawierająca próbkę dźwięku, jest dzielona na sekcje o długości (typowo) od 10 ms do 100 ms, które następnie są odtwarzane w innej kolejności, niż pierwotna, co pozwala na tworzenie nowych brzmień.

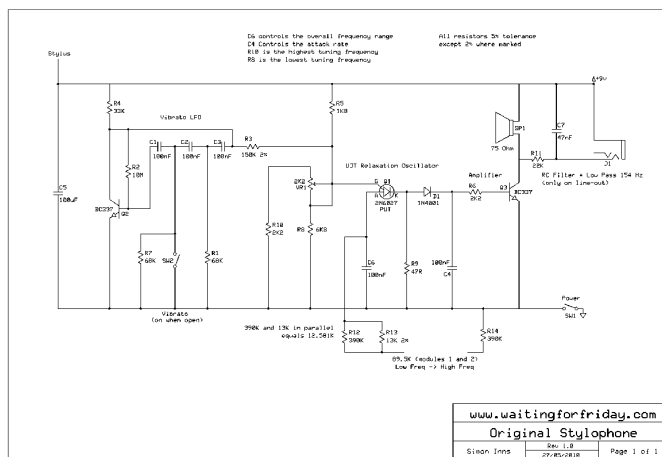
Synteza przez zniekształcanie fazy (phase distortion), opracowana przez firmę Casio, jest podobna do syntezy FM i wavetable – modulator zmienia chwilową fazę sygnału nośnego, najczęściej poprzez manipulację adresem tablicy z wartościami sygnału. Powstały przebieg wyjściowy przypomina kształtem sygnał syntezatora granularnego, jeśli użyć przebiegów prostych – takich jak sinus, przebieg trójkątny czy piłokształtny.

Anatomia syntezatora

Każdy syntezator zawiera trzy podstawowe elementy:

- oscylator,
- wzmacniacz,
- układu kontroli częstotliwości.

Na podstawie fotografii 1...3 Czytelnik mógłby dojść do wniosku, że jedyną metodą kontroli syntezatora jest zwykła klawiatura muzyczna, taka jak w fortepianie czy organach. Nie zawsze jednak jest to prawda. Prosty, by nie powiedzieć prymitywny, syntezatorem jest Stylophone (**fotografia 4**), instrument oparty na prostym oscylatorze relaksacyjnym. Rysik zamyka obwód oscylatora przez moduł składający się z zestawu rezystorów połączonych szeregowo, między którymi znajdują się odczepy w formie pól stykowych



Rysunek 1. Schemat oryginalnego minisyntezatora Stylophone z 1968 roku (<https://t.ly/TWybg>)

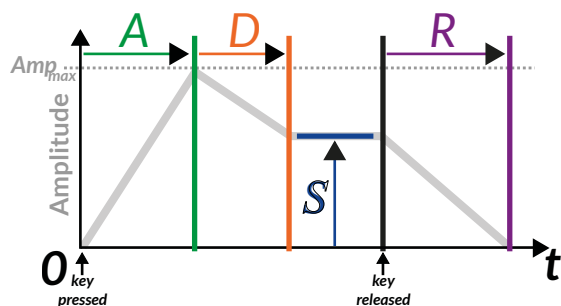
na płytce drukowanej, do których dotyka koniec rysika. **Rysunek 1** pokazuje schemat takiego układu. Wartości rezystorów są uszeregowane następująco:

- moduł 1: 21 kΩ, 2,05 kΩ, 2,18 kΩ, 2,32 kΩ, 2,45 kΩ, 2,59 kΩ, 2,75 kΩ, 2,91 kΩ, 3,08 kΩ, 3,27 kΩ;
- moduł 2: 3,46 kΩ, 3,66 kΩ, 3,88 kΩ, 4,11 kΩ, 4,35 kΩ, 4,59 kΩ, 4,84 kΩ, 5,15 kΩ, 5,49 kΩ, 5,85 kΩ.

Wariacje na temat tego instrumentu były realizowane z użyciem słynnego NE555 w roli oscylatora, a rysik wybierał rezystor ustalający częstotliwość drgań. Co ciekawe, jeden z wczesnych syntezatorów – Trautonium (**fotografia 5**) – również używał oscylatora



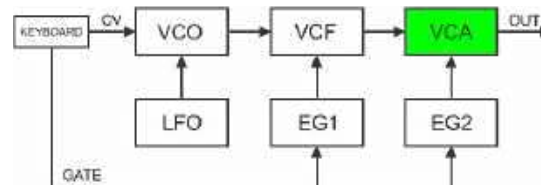
Fotografia 5. Syntezator Mixture Trautonium (<https://t.ly/Yck2p>)



Rysunek 2. Obwiednia ADSR (<https://t.ly/9jyqw>)

relaksacyjnego i rezystora, stworzonego ze struny owiniętej drutem oporowym i metalowej płyty, stanowiącej środkowy kontakt potencjometru. Dociskając strunę muzyk ustalał częstotliwość oscylatora, a drugą ręką kontrolował głośność/ekspresję. Z kolei inny instrument, Theremin, używał dwóch anten, do których użytkownik zbliżał ręce. Mniejsza antena kontrolowała głośność, zaś większa, pionowa, tworzyła obwód rezonansowy z ciałem muzyka. Poprzez ruch ręki operator mógł modyfikować częstotliwość tego obwodu, a różnica między jego sygnałem wyjściowym, a przebiegiem z drugiego generatora, stanowiła sygnał wyjściowy całego instrumentu.

Późniejsze syntezatory, zwłaszcza te produkowane począwszy od lat 70., zawierały już powszechnie klawiatury sterujące. Zamiast pojedynczego oscylatora, oferującego przebieg jednego typu, występowały rozbudowane bloki oscylatorów, z możliwością wyboru kształtu fali i jej częstotliwości. Dodatkowo pojawiły się też oscylatory LFO, generujące niskie częstotliwości w zakresie od 0,5 Hz do 10...20 Hz, dodatkowo modulujące główne oscylatory. Zrezygnowano też z ręcznej regulacji głośności, na rzecz generatora obwiedni zwanego też generatorem ADSR. Obwód ten, wyzwalany sygnałem sterującym z klawiatury (lub innego źródła), wytwarza sygnał pokazany na **rysunku 2**. Pierwsza część, A czyli Attack, określa szybkość narastania napięcia, a zarazem głośności sygnału z syntezatora – od zera do maksymalnej wartości. Sekcja D (Decay), określa czas opadania sygnału do wartości pośredniej. Fragment S (Sustain) definiuje pośrednią amplitudę sygnału, natomiast odcinek R (Release) określa czas opadania sygnału do zera po zwolnieniu klawisza. Niekiedy sekcja S może mieć ustaloną wartość czasu trwania, niezależną od stanu klawisza. Zwykle generator obwiedni i wzmacniacz sterowany jego przebiegiem wyjściowym znajdują się



Rysunek 3. Schemat blokowy prostego syntezatora (<https://t.ly/BsseH>)

na końcu ścieżki sygnałowej syntezatora. Wcześniej może znajdować się sekcja filtrów sterowanych napięciem, które dodatkowo mogą być modulowane oscylatorem LFO lub/i własnym generatorem obwiedni ADSR.

Rysunek 3 pokazuje schemat blokowy przykładowego syntezatora analogowego – sygnał CV to napięcie sterujące częstotliwością oscylatorów VCO, standardowo mieszczące się w zakresie 0...10 V (przeważnie przyjmuje się zależność 1 V/oktawę). Sygnał wyzwalający (Gate) kontroluje różne elementy syntezatora, w tym wypadku generatory obwiedni oznaczone literami EGx (Envelope Generator). Zaprezentowany na rysunku 3 podstawowy schemat blokowy prostego syntezatora nie wyczerpuje oczywiście wszystkich możliwości konstrukcyjnych. W praktyce syntezatory mogą mieć po kilka oscylatorów sterowanych wspólnym sygnałem CV, niezależnie filtry dla każdego z nich, a także mikser tych sygnałów, obwód modulacji czy dodatkowe filtry wyjściowe. Syntezatory modularne mogą zawierać wiele różnych typów oscylatorów, filtrów i wzmacniaczy sterowanych napięciem, a także innych modułów, a każdy z nich może mieć zarówno wejścia CV i Gate, jak i wyjścia tych sygnałów. Syntezatory monolityczne zwykle ograniczają się do dwóch lub trzech oscylatorów, jednego zestawu filtrów, jednego generatora LFO i jednego lub dwóch generatorów obwiedni. Często dodane są też efekty: pogłos (reverb) lub echo.

Zakończenie

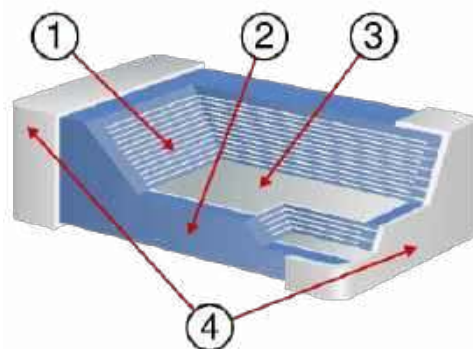
W następnej części przyjrzymy się wczesnym syntezatorom z okresu międzywojennego oraz instrumentom „syntezatoropodobnym”. W kolejnych odcinkach omówimy dokładniej poszczególne bloki syntezatora i zapoznamy się z ich zasadami działania oraz projektowania. Uzbrojony w tę wiedzę Czytelnik będzie mógł pokusić się o samodzielne zaprojektowanie i zbudowanie pełnoprawnego syntezatora analogowego.

Paweł Kowalczyk, EP

REKLAMA

Wszystko, co chcieliście wiedzieć o MLCC, ale baliście się zapytać

Elementy pasywne montowane powierzchniowo (SMD) – w szczególności rezystory i kondensatory – są wszechobecne w elektronice i choć często traktowane jako oczywiste elementy każdego układu, kryją w sobie wiele ciekawych aspektów konstrukcyjnych i eksploatacyjnych. W niniejszym artykule przyjrzymy się bliżej popularnym, wielowarstwowym kondensatorom ceramicznym MLCC (ang. Multilayer Ceramic Chip Capacitors). Omówimy stosowane w nich materiały i technologie, porównamy ich parametry (także te rzadziej eksponowane w notach katalogowych) oraz wskazówki praktyczne istotne dla projektantów i technologów odpowiedzialnych za produkcję urządzeń – zwłaszcza w przypadku branż krytycznych.



Rysunek 1. Budowa kondensatora wielowarstwowego (MLCC). 1 – warstwa dielektryczna, 2 – zewnętrzny korpus ceramiczny, 3 – warstwa przewodząca (okładka), 4 – końcówki lutownicze (<https://t.ly/QcMpX>)

Kondensator MLCC zbudowany jest z wielu naprzemiennie ułożonych warstw materiału ceramicznego (dielektryka) i wewnętrznych elektrod metalowych, które wyprowadzone są do zewnętrznych końcówek lutowniczych kondensatora. Dzięki takiej budowie uzyskuje się dużą pojemność w małej obudowie – kilkaset, a nawet tysiąc warstw może tworzyć jedną strukturę o wymiarach zewnętrznych rzędu ułamków milimetra. Na **rysunku 1** pokazano schematycznie budowę typowego kondensatora wielowarstwowego.

Rodzaje dielektryków – wpływ temperatury

Z materiałowego punktu widzenia wyróżnia się tzw. klasy dielektryków ceramicznych. Kondensatory **klasy I** (np. C0G/NP0) bazują na materiałach o bardzo stabilnych parametrach, przy czym najczęściej stosowanym izolatorem jest CaZrO_3 (cyrkonian wapnia). Minimalne zmiany pojemności pod wpływem temperatury czy pola elektrycznego (w wyniku braku domen ferroelektrycznych – dielektryki klasy I są paraelektrykami) sprawiają, że kondensatory te nie wprowadzają do sygnałów przemiennych zniekształceń spowodowanych histerezą. Mało tego – kolejną zaletą dielektryków klasy I jest praktycznie zerowy wpływ starzenia na wartość pojemności.

Czy są to zatem elementy idealne? Z punktu widzenia stabilności – jak najbardziej można je uznać za bliskie perfekcji. Niestety, dielektryki klasy I mają znikomą (w porównaniu do innych materiałów) przenikalność dielektryczną, przez co ich maksymalna pojemność jest mocno ograniczona. W katalogu jednego z największych na świecie dostawców komponentów elektronicznych można wprawdzie znaleźć kondensatory MLCC typu C0G o pojemności np. 470 nF, jednak ich rozmiar (wyrażony w systemie calowym) to aż... 2220 (długość 5,7 mm, szerokość 5,0 mm). Cena takich elementów także nie jest niska, zwłaszcza w porównaniu do typowych, tanich kondensatorów, które przy większych zakupach kosztują co najwyżej kilka groszy – tutaj trzeba się bowiem liczyć z cenami rzędu 10...30 zł/sztukę i to przy zamówieniach hurtowych. Na szczęście w praktyce dość rzadko zdarzają się sytuacje, by ultrastabilny kondensator o tak wysokiej pojemności rzeczywiście był potrzebny – przeważnie kondensatory z dielektrykiem klasy I są stosowane w obwodach radiowych (dopasowanie impedancji, filtry, sprzęganie), a także np. w oscylatorach kwarcowych.

Elementy o pojemności rzędu pikofaradów są na szczęście nieporównanie tańsze i w dodatku łatwo dostępne, także w małych rozmiarach (w tym 0402).

Warto w tym momencie powiedzieć nieco więcej o sposobie kodowania parametrów dielektryków klasy I. W **tabeli 1** zebrano informacje niezbędne do zdekodowania 3-pozycyjnego oznaczenia parametrów stabilności termicznej takiego dielektryka. Pierwsza litera definiuje liczbę znaczącą, niezbędną do obliczenia współczynnika temperaturowego pojemności zaś następująca po niej cyfra koduje mnożnik. Na trzeciej pozycji znajduje się znów oznaczenie literowe, tym razem jednak kodujące tolerancję współczynnika temperaturowego, wyrażoną w ppm/°C. I tak najpopularniejszy kod C0G oznacza (praktycznie) zerowy współczynnik termiczny, z tolerancją dopuszczalną na poziomie 30 ppm/°C (w całym zakresie temperatur). Z kolei kondensator z dielektrykiem U2J ma już znacznie mniejszą stabilność termiczną, gdyż jego współczynnik termiczny to aż 750 ppm/°C (–100/7,5) z tolerancją równą ± 120 ppm/°C – gorsze parametry są ceną za 2...4-krotnie większą pojemność, w porównaniu do kondensatora C0G o tych samych rozmiarach.

Kondensatory z dielektrykiem **klasy II** są zdecydowanie najliczniejsze w większości zastosowań (za wyjątkiem obwodów RF i innych

Tabela 1. Sposób oznaczania dielektryków klasy I (<https://t.ly/pGmWU>)

Pierwszy znak		Drugi znak		Trzeci znak	
Litera	Liczba znacząca współczynnika temperaturowego [ppm/°C]	Cyfra	Mnożnik	Litera	Tolerancja współczynnika temperaturowego [$\pm$ ppm/°C]
C	0,0	0	–1	G	± 30
B	0,3	1	–10	H	± 60
L	0,8	2	–100	J	± 120
A	0,9	3	–1000	K	± 250
M	1,0	4	+1	L	± 500
P	1,5	6	+10	M	± 1000
R	2,2	7	+100	N	± 2500
S	3,3	8	+1000		
T	4,7				
V	5,6				
U	7,5				

precyzyjnych układów, zdominowanych przez klasę I). Grupa materiałów klasy II obejmuje m. in. dielektryki X5R, X6R i X7R. Ceramika bazuje na ferroelektrycznym tytanianie baru, który pozwala uzyskać bardzo wysoką przenikalność dielektryczną (a więc dużą pojemność w przeliczeniu na jednostkę objętości), kosztem zdecydowanie gorszej stabilności parametrów. Z uwagi na diametralnie inny zakres zmienności pojemności w funkcji temperatury (wynikający z budowy dielektryka na poziomie mikroskopowym), sposób kodowania literowo-cyfrowego jest inny, choć także oparty na 3-znakowym kodzie. Sposób odczytywania kodów można znaleźć w tabeli 2 – jak widać, tutaj także całość dotyczy termicznej stabilności pojemności, ale nie mamy już bezpośredniej informacji o współczynniku temperaturowym. Zamiast niej poszczególne pozycje kodują:

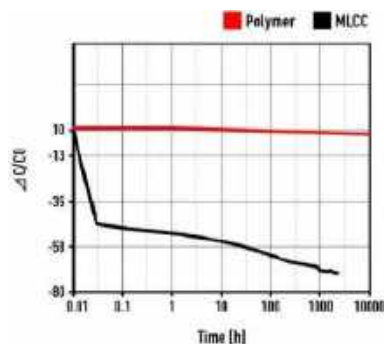
1. dolną granicę dopuszczalnego zakresu temperatur pracy,
2. górną granicę dopuszczalnego zakresu temperatur pracy,
3. zakres zmian pojemności w funkcji temperatury (w odniesieniu do granic określonych przez literę i cyfrę na początku oznaczenia).

Najpopularniejsze kondensatory X7R mogą zatem pracować w zakresie temperatur otoczenia od -55 do $+125^{\circ}\text{C}$, ale ich pojemność może zmieniać się w tych warunkach o $\pm 15\%$ wartości nominalnej.

W tym miejscu należy dodać, że istnieje także **klasa III**, również obejmująca materiały ferroelektryczne, ale tym razem już o ekstremalnie wysokiej przenikalności. Kondensatory wykonane z takich dielektryków wykazują bardzo duże zmiany pojemności, dochodzące nawet do -82% (w przypadku dielektryków z oznaczeniem V na trzeciej pozycji). Choć stabilność temperaturowa tych komponentów jest wręcz fatalna, to nie należy ich całkowicie deprecjonować – mogą okazać się niezastąpione w urządzeniach, które wymagają silnej miniaturyzacji, ale ponieważ zwykle pracują w warunkach zbliżonych do temperatury pokojowej, to nie trzeba przejmować się znanym spadkiem pojemności mierzonym w warunkach skrajnych. Rzecz jasna w urządzeniach przemysłowych, motoryzacyjnych czy wojskowych, które są narażone na zmiany temperatury w szerokim zakresie, zdecydowanie warto postawić na stabilniejsze kondensatory klasy II.

Zmiana pojemności w funkcji napięcia

Ważną cechą kondensatorów ceramicznych klasy II i III jest silna zależność pojemności od przyłożonego napięcia stałego (DC). Zjawisko to, nazywane efektem *DC bias*, powoduje, że rzeczywista pojemność kondensatora w układzie może być wielokrotnie niższa od nominalnej. Niestety, kod typu dielektryka (np. X7R) nie zawiera żadnej informacji o podatności na ten efekt – projektant musi zaglądać do charakterystyk zawartych w notach katalogowych. Ogólnie rzecz biorąc: im większa jest pojemność uzyskana w małym



Rysunek 2. Spadek pojemności pod napięciem, wykreślony w funkcji czasu (<https://t.ly/LkynT>)

Tabela 2. Sposób oznaczania dielektryków klasy II i III (wyróżnione gwiazdką i pogrubieniem)

Zródło: (<https://t.ly/pGMWU>)

Pierwszy znak		Drugi znak		Trzeci znak	
Litera	Dolny limit zakresu temperatur $^{\circ}\text{C}$	Cyfra	Górny limit zakresu temperatur $^{\circ}\text{C}$	Litera	Zakres zmian pojemności w funkcji temperatury [%]
X	-55	2	$+45$	D	$\pm 3,3$
Y	-30	4	$+65$	E	$\pm 4,7$
Z	$+10$	5	$+85$	F	$\pm 7,5$
		6	$+105$	P	± 10
		7	$+125$	R	± 15
				S	± 22
				T*	$+22/-33$
				U*	$+22/-56$
				V*	$+22/-82$

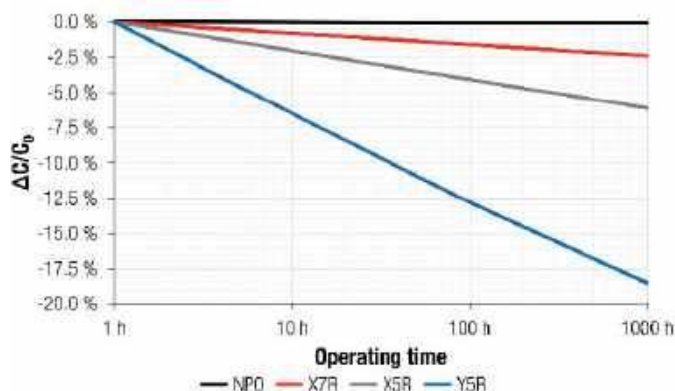
kondensatorze (czyli im cieńsze warstwy dielektryka i im więcej warstw), tym silniejszy spadek pojemności pod wpływem pola elektrycznego. Dokładniejszy opis wpływu efektów ferroelektrycznych na histerezę krzywej pojemności w funkcji napięcia można znaleźć w artykule „Komponenty bierne pod lupą” – nie będziemy więc ponownie zagłębiać się w tę tematykę.

Trzeba natomiast zdecydowanie zwrócić uwagę na inne zjawisko związane z ferroelektrycznym charakterem dielektryków klasy II i III. Okazuje się bowiem, że **spadek pojemności nie następuje natychmiast** po przyłożeniu napięcia – część wolniejszych domen ustawia się z pewnym opóźnieniem. Dlatego po załączeniu napięcia obserwuje się początkowy szybki spadek pojemności, a następnie wolniejszy, dodatkowy ubytek postępuje w ciągu kolejnych minut i godzin pracy (rysunek 2). W układach pracujących z przebiegami zmiennymi regularne wahania napięcia „resetują” część uwieczonych dipoli. Natomiast w zastosowaniach takich, jak chociażby odsprężanie i filtracja zasilania, inżynierska praktyka nakazuje zaplanować użycie kondensatora X7R lub X5R z dużym nadatkiem pojemności, jeśli ma on pracować z relatywnie wysokim (względem nominalnego) napięciem stałym. Pomocne bywa wybranie elementu w większej obudowie i/lub o znacznie wyższym napięciu nominalnym niż rzeczywiste napięcie robocze – wtedy grubsze warstwy dielektryka i słabsze natężenie pola (w przeliczeniu na warstwę) pozwolą zredukować efekt *DC bias*.

Effekt starzenia dielektryka

W tym momencie należy wspomnieć o kolejnym problemie, jakim jest zjawisko stopniowego spadku pojemności w czasie od momentu wyprodukowania lub od ostatniego wygrzania kondensatora. Jest to tzw. starzenie dielektryka (tzw. *aging*), przebiegające (z grubsza) wykładniczo – pojemność maleje o pewien procent na dekadę czasu (np. 1 godzina $\rightarrow$ 10 godzin $\rightarrow$ 100 godzin itd.) – patrz rysunek 3. Typowe wartości to ok. 2...5% na dekadę czasu dla (odpowiednio) X7R oraz X5R. Oznacza to, że po dłuższym czasie pojemność kondensatora może spaść nawet o 40...50% pierwotnej wartości! Co ważne, podgrzanie kondensatora powyżej tzw. temperatury Curie (około $120...130^{\circ}\text{C}$ w przypadku BaTiO_3) przywraca mu pełną pojemność – struktura domen dipolowych reorganizuje się (dla celów standaryzacji określa się więc pojemność poprzez pomiar 24 godziny po wygrzaniu). Starzenie nie jest przy tym traktowane jako uszkodzenie elementu, lecz stanowi nieuniknioną właściwość samego materiału. W przypadku kondensatorów klasy I zjawisko starzenia praktycznie nie występuje, a przynajmniej jego wpływ jest pomijalny w porównaniu do pozostałych klas.

Trzeba przy tym pamiętać, że starzenie jest zjawiskiem niezależnym od opisanego wcześniej spadku pojemności pod napięciem



Rysunek 3. Porównanie efektu starzenia dla czterech różnych dielektryków (<https://t.ly/MbtE5>)

– jest ono bowiem definiowane jako „samoczynna” zmiana pojemności w warunkach temperatury pokojowej (np. 20°C) i **przy zerowym napięciu**. Jeżeli dodamy do tego dość sporą tolerancję pojemności (klasyczne rozrzuty produkcyjne rzędu $\pm 10...20\%$) oraz opisane wcześniej zależności pojemności od temperatury i napięcia, to otrzymamy naprawdę złożony obraz sprawy – jak widać na **rysunku 4**, kondensatorom klasy II (nie wspominając o III) zdecydowanie nie można zbyt ufać pod względem pojemności nominalnej. Wartość ta jest raczej mocno orientacyjna, ale pamiętając o wspomnianych w artykule efektach można zdecydowanie bardziej świadomie stosować te elementy w niemal dowolnych urządzeniach elektronicznych.

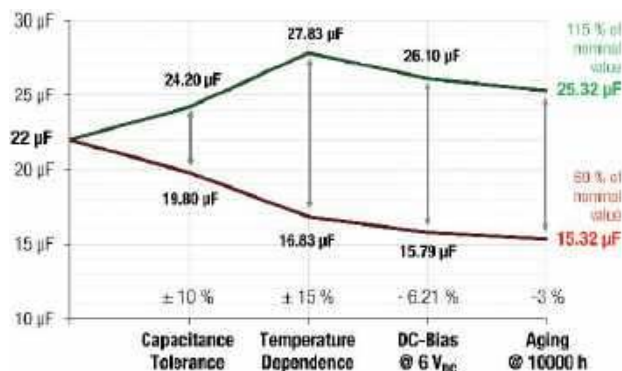
Inne czynniki wpływające na pojemność

Rzeczywista pojemność kondensatorów MLCC zależy nie tylko od wymienionych powyżej czynników – podlega ona także (choć w znacznie mniejszym stopniu) działaniu innych efektów fizycznych zachodzących w dielektryku. Przejrzenie dokumentacji „pierwszego z brzegu” kondensatora do aplikacji motoryzacyjnych ujawnia szereg kolejnych parametrów środowiskowych, wraz z dopuszczalnymi odchyłkami pojemności powodowanymi przez zmiany tychże wielkości. Przykład? Proszę bardzo. Firma Samsung wskazuje w dokumentacji kondensatora CL10C101JC81PNC (100 pF, C0G, 100 V, 0603, AEC-Q200), że pojemność nie powinna zmieniać się o więcej, niż 2,5% (lub 0,25 pF – w zależności od tego, która wartość jest większa) pod wpływem działania wilgoci na poziomie 85% RH i w temperaturze 85°C. Te same granice dotyczą także wpływu wstrząsów, wibracji, temperatury w procesie lutowania, „strzałów” ESD czy wreszcie... naprężeń, którym poddawane są wyprowadzenia lutownicze kondensatora. 5-procentowa odchyłka jest natomiast dopuszczalna w teście na zginanie płytki drukowanej, na której zamontowany jest kondensator.

Nie tylko pojemność, czyli co nieco o rezystancji dielektryka

Jak można wywnioskować z dotychczasowego opisu, o parametrach kondensatorów MLCC (i nie tylko) zdecydowanie w największym stopniu decyduje sam dielektryk. A, jak powszechnie wiadomo, rolą dielektryka jest nie tylko mechaniczne rozdzielanie okładek kondensatora, ale przede wszystkim wprowadzenie pomiędzy nie pewnej izolacji elektrycznej. Natomiast skoro mowa o izolacji – to czego można się spodziewać odnośnie jej rezystancji?

Okazuje się, że na ten parametr wpływają te same czynniki, które działają także na pojemność



Rysunek 4. Ilustracja wpływu różnych efektów niepożądanych, zachodzących w dielektryku, na wypadkową pojemność kondensatorów MLCC (<https://t.ly/MbtE5>)

– choć w zupełnie innych granicach. Posłużymy się w tym miejscu przykładem tego samego kondensatora marki Samsung (CL10C101JC81PNC), którego parametrom przyjrzelśmy się w poprzedniej części artykułu. Typowa wartość rezystancji izolacji w warunkach nominalnych wynosi wprawdzie dość sporo, bo minimum 100 GΩ (lub 1 GΩ·µF), jednak w większości rygorystycznych testów (np. odporności na lutowanie, cykle termiczne, wibracje, udary mechaniczne czy wyładowania ESD) wartości progowe są 10-krotnie niższe (10 GΩ) – patrz **rysunek 5**. Zdecydowanie najmniejsza jest dopuszczalna rezystancja izolacji w warunkach zwiększonej wilgotności – tutaj wymagania testu spadają do zaledwie 500 MΩ.

Z czego to wynika? Jak nietrudno się domyślić, porowata struktura ceramiki dielektryka jest w pewnym stopniu higroskopijna. Dlatego właśnie kondensatory MLCC „nie lubią” warunków o dużej wilgotności, gdyż wpływa ona nie tylko na zmniejszenie pojemności i rezystancji oraz zwiększenie wartości prądów upływu, ale potencjalnie może doprowadzić nawet do mechanicznych uszkodzeń struktury kondensatora (zwłaszcza przy udziale dodatkowych narażeń, np. wibracji, wstrząsów czy wahań temperatury).

Wpływ warunków montażu i naprężeń

Kondensatory ceramiczne są kruche – zbudowane z twardej ceramiki, przez co podatne na pęknięcia. Jedną z najczęstszych przyczyn awarii MLCC jest tzw. *flex cracking*, czyli pęknięcie na skutek wygięcia płytki drukowanej. Odształcenie laminatu może nastąpić np. przy dociskaniu pobliskiego złącza, w czasie przykręcania PCB do obudowy, w wyniku upadku urządzenia, a w skrajnych przypadkach nawet z powodu różnic rozszerzalności cieplnej (naprężenia termiczne spowodowane gradientem temperatury). Pęknięcie kondensatora najczęściej rozpoczyna się przy krawędzi

C. Reliability Test and Judgement Condition		
Test Items	Performance	Test Condition
High Temperature Exposure	Appearance: No abnormal exterior appearance Capacitance Change: Within 2.5% or 0.25pF Whichever is larger Q: 1000min. IR: More than 1000MΩ or 500MΩ·µF Whichever is smaller	Unpowered, 1000 hours @ Max. Temperature Measurement at 24±2hrs after test conclusion
Temperature Cycling	Appearance: No abnormal exterior appearance Capacitance Change: Within 2.5% or 0.25pF Whichever is larger Q: 1000min. IR: More than 1000MΩ or 500MΩ·µF Whichever is smaller	1000 Cycles Measurement at 24±2hrs after test conclusion 1 cycle condition: -55±0/-31(30s/3min) → Room Temp. (1min) → 125±3/-0°C(30s/3min) → Room Temp. (1min)
Destructive Physical Analysis	No Defects or abnormalities	Per EIA 489
Humidity Bias	Appearance: No abnormal exterior appearance Capacitance Change: Within 2.5% or 0.25pF Whichever is larger Q: 200min. IR: More than 500MΩ or 25MΩ·µF Whichever is smaller	1000 hours 85°C/85% RH, Rated Voltage and 1.3~1.5V, Add 100kΩ resistor The charge/discharge current is less than 50µA
High Temperature Operating Life	Appearance: No abnormal exterior appearance Capacitance Change: Within 3% or 0.3pF Whichever is larger Q: 350min. IR: More than 1000MΩ or 50MΩ·µF Whichever is smaller	1000 hours @ 125°C, 200% Rated Voltage, Measurement at 24±2hrs after test conclusion The charge/discharge current is less than 50µA

Rysunek 5. Fragment dokumentacji kondensatora CL10C101JC81PNC (<https://t.ly/cEDXd>)



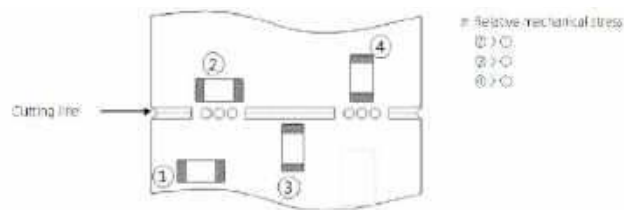
Fotografia 1. Skośnie przebiegająca szczelina będąca wynikiem uszkodzenia typu flex cracking (<https://t.ly/s67CV>)

jego wyprowadzenia i przebiega skośnie przez warstwy do wnętrza elementu. Jeśli trasa pęknięcia odsłoni dwie sąsiednie elektrody wewnętrzne o przeciwnej biegunowości, to kondensator ulegnie zwarceniu lub przynajmniej znacząco zwiększy się jego upływność. A to już recepta na bardzo niebezpieczny tryb uszkodzenia – wewnętrzne zwarcie MLCC może spowodować poważną awarię obwodów zasilania. **Fotografia 1** pokazuje przykład takiego pęknięcia spowodowanego wygięciem płytki – widoczna szczelina rozdziela warstwy dielektryka. Rozwiązaniem problemu jest daleko posunięta ostrożność w zakresie projektowania PCB oraz obudowy urządzenia, sposobu mocowania płytki, a nawet doboru jej geometrii (stosunku wymiarów) – wszystko zależy oczywiście od konkretnej sytuacji projektowej. Ważne jest maksymalne ograniczenie ugięć (np. montaż kondensatorów równoległe do spodziewanej osi zginania PCB) oraz mocowanie elementów w miarę możliwości daleko od krytycznych punktów PCB, w tym otworów montażowych, złączy czy perforacji służących do wyłamania płytki z panelu (tzw. *mouse bites*). Przykładowe zalecenia można zobaczyć na **rysunkach 6...9**.

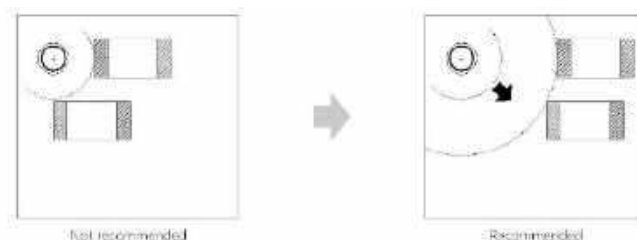
Częściowym rozwiązaniem problemu jest także zastosowanie kondensatorów o tzw. elastycznych wyprowadzeniach (*soft termination*). Elementy te mają dodatkową warstwę materiału (zwykle żywicy epoksydowej z cząstkami metalu) między ceramicznym korpusem a zewnętrzną metalizacją elektrod (**rysunek 10**). Warstwa ta działa jak „bufor” odształceń – kondensator wyposażony w elastyczne końcówki nie pęknie od razu, gdyż epoksyd skutecznie zniweluje mniejsze naprężenia. W skrajnym przypadku, czyli w razie nadmiernego wygięcia płytki, po prostu zerwą się połączenia wewnętrzne kondensatora – przez co utraci on swoją pojemność (uszkodzenie typu *open*), ale nie doprowadzi do (w większości przypadków) znacznie groźniejszego zwarcia.

Niestety prezentowane rozwiązanie ma także swoją cenę w postaci zwiększonej rezystancji połączeń kondensatora (warstwa epoksydowa przewodzi prąd gorzej, niż czysty metal), co zmniejsza efektywność filtracji lub odsprężania zasilania. Warto pamiętać o tym kompromisie i w razie potrzeby zastosować dodatkowe środki zapobiegawcze (np. poprzez zwiększenie liczby kondensatorów równoległych czy też zastosowanie dodatkowego kondensatora elektrolytycznego o niskiej impedancji).

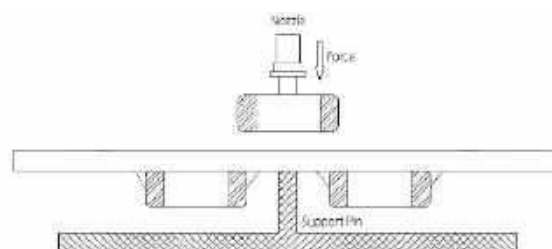
W tym miejscu warto dodać, że technologia *soft termination* to niejedyna modyfikacja klasycznej konstrukcji MLCC. Na rynku dostępne są także kondensatory określane mianem *floating electrode*, co oznacza dosłownie „pływającą elektrodę” (**rysunek 11**). Trik jest prosty – żadna z okładek nie przechodzi z jednej strony kondensatora (gdzie łączy się z wyprowadzeniem lutowniczym) na drugą (tj. w pobliżu przeciwnego wyprowadzenia). Okładki główne (połączone z wyprowadzeniami) są bowiem znacznie krótsze niż w typowych kondensatorach, gdyż nie muszą zazębiać się z ze sobą – współpracują z blokiem niepodłączonych pól przewodzących, znajdujących się w środku korpusu kondensatora. W ten sposób całość tworzy de facto dwa kondensatory połączone szeregowo. Opisywana konstrukcja, jakkolwiek niezbyt optymalna (wypadkowa pojemność jest



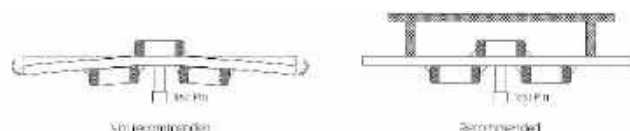
Rysunek 6. Przykład prawidłowego (1) i nieprawidłowych (2, 3, 4) pozycji kondensatorów MLCC względem perforacji służących do depanelizacji PCB (<https://t.ly/kwPHT>)



Rysunek 7. Zalecenie dotyczące umieszczania kondensatorów MLCC (w miarę możliwości) z dala od otworów montażowych PCB (<https://t.ly/kwPHT>)



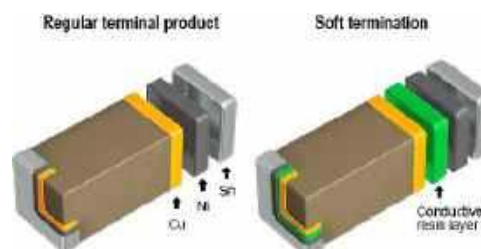
Rysunek 8. W przypadku montażu dwustronnego należy unikać nadmiernych naprężeń PCB podczas układania komponentów za pomocą automatu pick & place. W celu minimalizacji ryzyka uszkodzeń należy zastosować dodatkowe punkty podparcia płytki (<https://t.ly/kwPHT>)



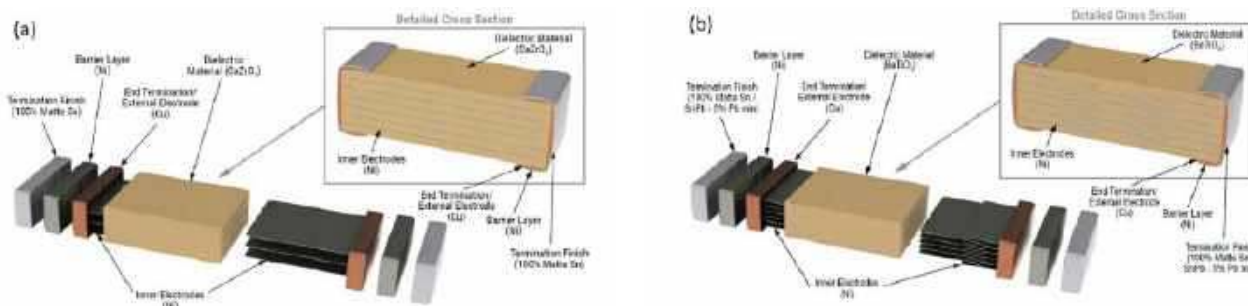
Rysunek 9. Te same zalecenia, które dotyczą montażu pick & place (patrz opis pod rysunkiem 8) można także zastosować do pinów testowych (<https://t.ly/kwPHT>)

dwukrotnie mniejsza niż w porównywalnym kondensatorze o standardowej budowie), ma niezwykle istotną zaletę: żadne uszkodzenie (w postaci pęknięcia termicznego lub mechanicznego) nie może doprowadzić do zwarcia obu wyprowadzeń kondensatora, gdyż zawsze pomiędzy elektrodami pływającymi, a głównymi okładkami, istnieje przynajmniej z jednej strony izolacja w postaci nieprzerwanych warstw dielektryka.

Połączenie technologii *floating electrode* z opisaną wcześniej *soft termination* pozwala konstruktorom uzyskać jeszcze wyższy poziom niezawodności. Seria FF-CAP marki Kemet integruje zalety obydwu



Rysunek 10. Budowa klasycznego kondensatora MLCC (po lewej) i kondensatora typu soft termination (po prawej) (<https://t.ly/c8DGv>)



Rysunek 11. Porównanie budowy klasycznego kondensatora MLCC (a) i kondensatora typu floating electrode (b). Źródło: <https://t.ly/OZkE4>

wspomnianych rozwiązań, wpasowując się w potrzeby branż krytycznych – w tym motoryzacyjnej, medycznej, kosmicznej, lotniczej, telekomunikacyjnej czy przemysłowej.

Tryby uszkodzeń kondensatorów

Kondensatory MLCC podlegają dwóm głównym rodzajom uszkodzeń, powodujących albo znaczne zwiększenie prądu upływu (spadek wypadkowej rezystancji izolacji, w skrajnym przypadku w postaci zwarcia) lub całkowitą (bądź częściową) utratę pojemności. Przyczyny awarii mogą być jednak diametralnie różne, nawet w obrębie jednej grupy uszkodzeń.

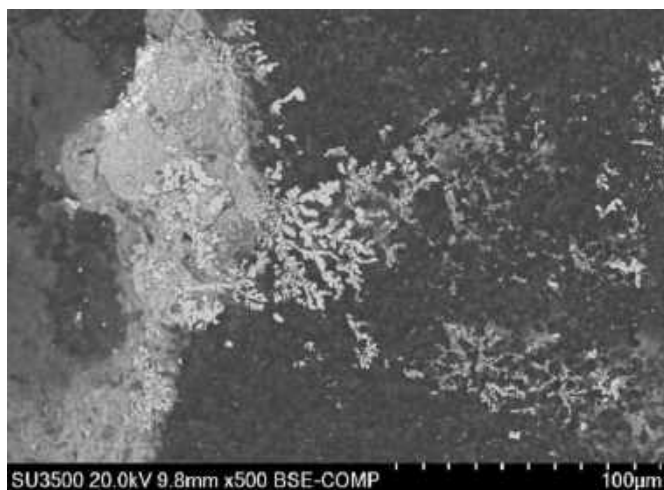
Jak już wspomnieliśmy, najczęstsze (i zwykle najpoważniejsze w skutkach) awarie MLCC to pęknięcia wewnętrzne, prowadzące do powstania przerwy lub zwarcia (w zależności od tego, czy dojdzie do kontaktu pomiędzy przeciwnymi okładkami). Są one przeważnie spowodowane naprężeniami mechanicznymi lub termicznymi. To jednak dopiero początek długiej listy potencjalnych źródeł awarii.

Przebite dielektryka wskutek przepięcia może doprowadzić do trwałego zwarcia na skutek zniszczenia cienkich warstw dielektryka pomiędzy okładkami pracującymi na przeciwnych potencjałach. Należy brać także pod uwagę potencjalne uszkodzenia spowodowane efektem migracji cząstek metalu (cyny, srebra lub miedzi), które – choć rzadziej spotykane – w obecności wilgoci i napięcia stałego mogą doprowadzić do zwiększenia prądu upływu, a nawet powstania niskooporowej ścieżki („małego zwarcia”). Przykład można zobaczyć na **fotografii 2**. W analizie konkretnego przypadku trzeba oczywiście wykluczyć udział czynników leżących de facto poza samym kondensatorem. Zewnętrzna ścieżka przewodząca może mieć np. postać kulki cyny, powstałej w wyniku błędów technologicznych w procesie lutowania rozpliwowego.

Innym rodzajem uszkodzenia jest wewnętrzna delaminacja bloku ceramicznego, spowodowana ukrytymi wadami pojawiającymi się na etapie produkcji kondensatora. Ten typ awarii ma tendencję do narastania w wyniku naprężeń termicznych i mechanicznych, powstających już na etapie eksploatacji i stanowi niejako „fabryczną słabość” danego komponentu, której rzecz jasna nie sposób wykryć na etapie montażu (chyba, że zostanie ona przypadkiem zauważona w wysokorozdzielczym obrazowaniu rentgenowskim). Przykład efektownej delaminacji, wywołanej w tym przypadku impulsem wysokiego napięcia, można zobaczyć na **fotografii 3**.

Trochę egzotyki

Szeroko opisane w niniejszym artykule problemy związane z podatnością kondensatorów MLCC na pękanie, doprowadziły producentów do prostego wniosku – doskonałą metodą na odizolowanie wrażliwej ceramiki od krytycznie silnych naprężeń mechanicznych może być... montaż elementu na metalowych „nogach”. Tak powstały kondensatory odporne na naprawdę trudne warunki pracy (**fotografia 4**) – w testach laboratoryjnych, wykonywanych na standardyzowanych płytkach drukowanych o ściśle określonej geometrii i wymiarach, elementy te są w stanie przetrwać ugięcie PCB rzędu 6 mm, podczas gdy dla typowych elementów granica odporności nie przekracza zwykle 3 mm.



Fotografia 2. Obraz powierzchni kondensatora MLCC zarejestrowany za pomocą mikroskopu skaningowego. Widoczne efekty elektromigracji srebra, która doprowadziła do powstania niskooporowej ścieżki zewnętrznej pomiędzy wyprowadzeniami lutowniczymi komponentu (<https://t.ly/j2QdL>)



Fotografia 3. Kondensator uszkodzony przez wysokie napięcie – widoczne poważne uszkodzenia zewnętrzne oraz obszerna delaminacja wewnątrz korpusu elementu (https://t.ly/Kul_1)

W ślad za tego rodzaju komponentami poszły także całe bloki kondensatorów, połączonych fabrycznie za pomocą wspólnych, metalowych blaszek. Takie równoległe baterie kondensatorowe występują zarówno w wersjach z dwoma lub trzema MLCC (**fotografia 5**), jak i w postaci rozbudowanych bloków złożonych z wielu, ciasno ułożonych kondensatorów (**fotografia 6**).

Co ciekawe, historia (a raczej miniaturyzacja) w pewnym sensie zatoczyła koło – rozwój komponentów takich, jak na **fotografii 6**, doprowadził do powstania nieco większych zestawów, w których blaszki łączące poszczególne MLCC mają wyprowadzone... klasyczne piny do montażu THT (sic!). Przykład takiego egzotycznego tworu można zobaczyć na **fotografii 7**. Jego pojemność wynosi 220 nF zaś napięcie robocze to 630 V. A cena? ponad 500 zł netto przy zamówieniach



Fotografia 4. Kondensator z metalowymi wyprowadzeniami zmniejszającymi naprężenia mechaniczne, przenoszone z płytki drukowanej na ceramiczny korpus elementu (<https://t.ly/u9W1Q>)



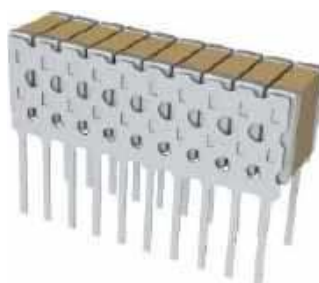
Fotografia 5. 2- oraz 3-elementowe bloki kondensatorów MLCC marki TDK (https://t.ly/_DIgW)



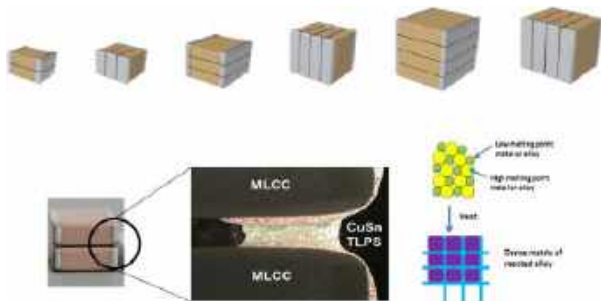
Fotografia 6. Rozbudowane baterie kondensatorów MLCC (<https://t.ly/Gjhd5>)

do 10 sztuk. Tak wysoki koszt wynika zarówno z doskonałego dielektryka (w tym przypadku jest to COG!), ale także z bardzo szerokiego zakresu temperatur pracy, dochodzącego do 200°C.

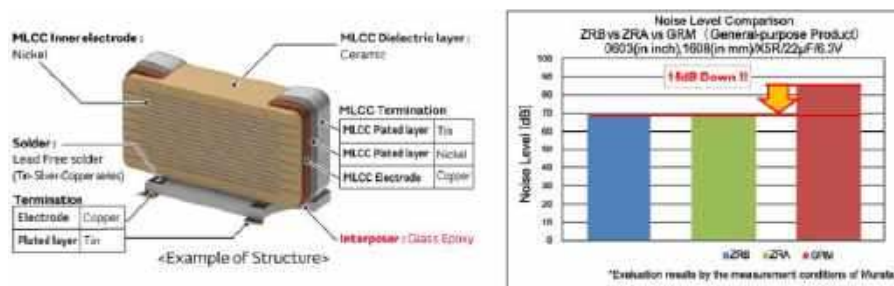
Jeszcze innym przykładem interesującego rozwiązania opartego na dość standardowych kondensatorach MLCC, jest technologia KONNEKT, opatentowana przez firmę Kemet. Poszczególne kondensatory w zestawie są ze sobą połączone bezpośrednio – zamiast blaszek (zgrzewanych lub lutowanych) połączenie mechaniczne i elektryczne jest wykonane przez zastosowanie specjalnej techniki określonej mianem TLPS (ang. *Transient Liquid Phase Sintering* – przejściowe spiekanie w fazie ciekłej) – patrz rysunek 12. Powstały w ten sposób



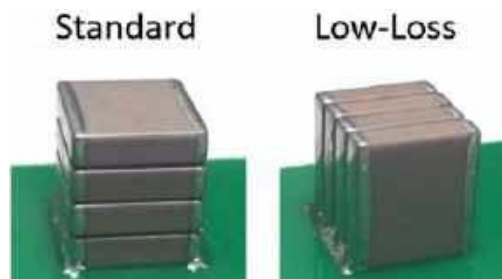
Fotografia 7. Bateria kondensatorów MLCC z wyprowadzeniami do montażu przewlekane (<https://t.ly/0FoyC>)



Rysunek 12. Technologia KONNEKT marki Kemet (<https://t.ly/9xaZg>)



Rysunek 13. Kondensator z serii ZRB marki Murata, przeznaczony do systemów wymagających cichej pracy przy napięciu szybkozmiennym (<https://t.ly/VEJ3l>)



Fotografia 8. Możliwe orientacje montażowe kondensatorów z serii KONNEKT – standardowa oraz niskosłupna (<https://t.ly/9xaZg>)

material trwale spaja sąsiadujące kondensatory, a w dodatku zapewnia doskonałą lutowalność. Główną zaletą takiego rozwiązania jest możliwość lutowania zestawu w dwóch orientacjach: standardowej (kondensatory są ustawione poziomo, jeden nad drugim) lub bocznej (kondensatory są ustawione pionowo obok siebie) – patrz fotografia 8. Zaletą drugiego rozwiązania jest znacznie niższa wartość ESR, gdyż każdy z kondensatorów składowych w ramach zestawu ma bezpośrednie połączenie z padami na PCB – genialne w swojej prostocie.

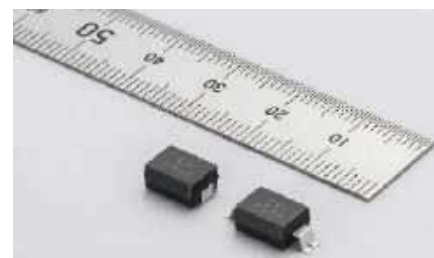
To jednak wciąż nie wszystko, na co stać najbardziej zaawansowanych producentów kondensatorów. Firma Murata wprowadziła na rynek kondensatory o niskim poziomie szumu akustycznego, montowane powierzchniowo za pomocą specjalnej przekładki (tzw. *interposer*), tłumiącej częściowo drgania powstające w strukturze kondensatora w wyniku efektu piezoelektrycznego. Takie rozwiązanie pozwala, zdaniem producenta, na uzyskanie o 15 dB cichszej pracy w porównaniu ze standardowymi modelami kondensatorów (rysunek 13).

Na koniec prezentacji kondensatorów „egzotycznych” pozostawiliśmy jeszcze jeden ciekawy twór. Umieszczenie kondensatora MLCC (o odpowiednich parametrach) w obudowie o konstrukcji znanej z typowych diod prostowniczych bądź transyli (fotografia 9) pozwoliło inżynierom marki Murata uzyskać elementy spełniające wymogi normy IEC60384-14 dla kondensatorów Y2, o odstępach powierzchniowych między wyprowadzeniami równym aż 10 mm. Takie rozwiązanie zostało opracowane z myślą o aplikacjach związanych z restrykcyjnymi wymogami bezpieczeństwa elektrycznego, m.in. w motoryzacji czy energoelektronice.

Podsumowanie

Jak widać, kondensatory MLCC wcale nie są prostymi i oczywistymi elementami – a przynajmniej nie w takim stopniu, jak widzi je wielu konstruktorów niezagłębiających się w zawile technika. Jeżeli jednak projektowane przez nas urządzenia muszą spełniać naprawdę rygorystyczne wymagania w zakresie niezawodności, bezpieczeństwa czy stabilności termicznej, warto zajrzeć w głąb dokumentacji udostępnianej przez producentów. Często znajdziemy tam bowiem wskazówki, które pozwolą uniknąć wielu niespodziewanych problemów, niezwykle trudnych do zdiagnozowania na późniejszych etapach cyklu życia produktu.

inż. Przemysław Musz, EP



Fotografia 9. Kondensatory zalewane żywicą, o zwiększonym odstępzie powietrznym i powierzchniowym (<https://t.ly/ms51p>)



Komponenty bierne pod lupą

Komponenty bierne, czyli rezystory, kondensatory i inne podzespoły, na których działanie z reguły mamy niewielki wpływ, zwykle nie są oczkiem w głowie projektantów elektroniki. Moc się zgadza, napięcie też, nie ma się czym więcej przejmować – ruszamy z produkcją. Tylko czy aby na pewno takie podejście zawsze jest właściwe?

Gdyby do rezystorów, kondensatorów i podobnych im podzespołów podchodzić na gruncie teorii obwodów – ale takiej czystej, wręcz purystycznej – mógłbym w tym miejscu zakończyć niniejszy artykuł wielką, efektowną kropką. Może nawet z wykrzyknikiem dla podkreślenia wagi tego faktu. Niestety po wyjściu poza zakres I semestru studiów okazuje się, że nawet najbardziej prymitywne elementy elektroniczne wcale takie prymitywne nie są. I nie chodzi bynajmniej o to, że to z nimi coś jest nie tak. One są w porządku, robią swoją robotę. Chodzi o to, że za każdym takim podzespołem stoi niemały kawał fizyki, który odpowiada za jego funkcjonowanie, a także model próbujący opisać te zjawiska w sposób zrozumiały dla elektronicznej braci, czyli... schemat zastępczy.

Schematy zastępcze też nie są sobie równe. Inaczej będziemy rozpatrywali rezystor w obwodzie wielkiej częstotliwości, w którym dochodzi kwestia jego indukcyjności i pojemności pól lutowniczych. Jeżeli dochodziłoby do jego silnego nagrzewania się, wówczas warto uwzględnić zmiany rezystancji wywołane tym faktem. Jeszcze inaczej będziemy analizować wpływ owego rezystora na działanie obwodu niskosumnego (a przynajmniej takiego, który powinien być niskosumny według założeń projektu). Dochodzą do tego inne zjawiska, z reguły rzadko wymieniane przez producentów, a które cechują jakiś rodzaj podzespołów, na przykład mikrofonowanie niektórych rodzajów kondensatorów.

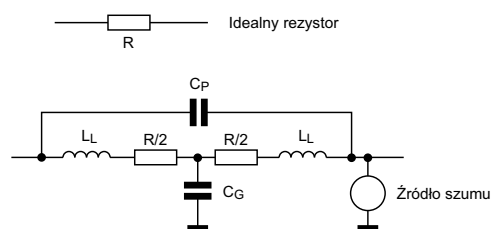
Myślę, że w 90% przypadków możemy zastosować podzespół pasujący wartością (rezystancją w przypadku rezystorów lub pojemnością w przypadku kondensatorów) i nie zastanawiać się nad nim

ani chwili dłużej. Rezystor 10 kΩ przy podciąganiu wyprowadzenia mikrokontrolera do linii zasilającej lub kondensator 100 nF przy odprężaniu zasilania prostego układu scalonego to rzeczy, nad którymi – w zdecydowanej większości przypadków – nikt się nie rozczula. Jednak nawet w relatywnie prostych układach bywają miejsca newralgiczne, czule na pewne cechy podzespołów biernych, zaś w szczególnie wymagających obszarach elektroniki (wojsko, kosmos, radary) całe projektowanie musi przebiegać w ściśle określony sposób.

W tym artykule postaram się poruszyć zagadnienia, które dotyczą pozostałej części (czyli około 10%) przypadków. Niekiedy ich pominięcie może poskutkować tym, że układ będzie się zachowywał w sposób zupełnie inny od zamierzonego: inne pasmo przenoszenia, dziwny szum na wyjściu, trzaski w sygnale, a nawet przedostające się do toru analogowego buczenie o częstotliwości 50 Hz, którego nie da się wyeliminować pomimo zastosowania ekranów elektromagnetycznych i wielu innych osłon.

Rezystory

To, co powinien robić rezystor, dobrze opisuje prawo Ohma. Lecz w rzeczywistości jest tak, że rezystory reprezentują sobą całą masę innych cech, niekoniecznie dla nas pożądaných. Przyjrzyjmy się tym zjawiskom, które występują najczęściej. Choć i tak rezystory należą do komponentów, które w znacznym stopniu są bliskie elementom idealnym, oczywiście w granicach prawidłowych warunków pracy.



Rysunek 1. Schemat zastępczy rezystora [1]

Na **rysunku 1** znajduje się dość powszechnie stosowany schemat zastępczy rezystora w zestawieniu z jego idealnym, „teoriobwodowym” modelem. Nietrudno rozszyfrować jego poszczególne składowe: R to rezystancja, LL to indukcyjność doprowadzeń, CP jest pojemnością między wyprowadzeniami zaś CG to pojemność do masy. Nie zapomniano też o źródle reprezentującym szum wprowadzany przez ten element.

Pojemność

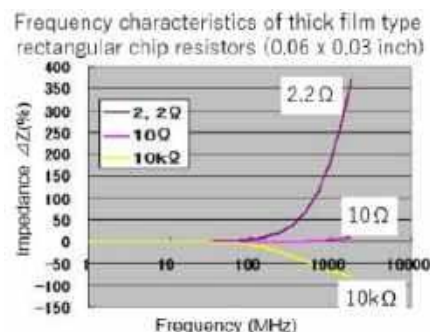
Pojemność między wyprowadzeniami zależy bezpośrednio od gabarytów danego rezystora. Według dokumentów Texas Instruments może ona wynosić od 0,01 pF do 0,5 pF [1], zależnie od wielkości obudowy. Kształt i ułożenie pól lutowniczych również wpływają na ten parametr i nie są zależne od samego elementu. Producenci nie chwalią się tym parametrem, bo jest trudny do zmierzenia i rzadko przydatny – choć warto mieć na uwadze jego istnienie.

Podobnie ma się sprawa z pojemnością do masy, która zależy już prawie wyłącznie od powierzchni pól lutowniczych i od konstrukcji stosu warstw płytki. Jeżeli zależy nam na dobrym poznaniu tej wartości, warto byłoby zaprzęgnąć do pracy symulator. Poza układami wielkiej częstotliwości, które w oczywisty sposób są czułe na tak niewielkie pojemności, również oscylatory (zwłaszcza w układzie Colpittsa) muszą mieć dobrze kontrolowane wszystkie pojemności pasywnicze.

Indukcyjność

Rezystor to odcinek przewodzącego materiału, więc sam z siebie cechuje się pewną indukcyjnością. Sprawy nie upraszcza fakt, że rezystory cienkowarstwowe przewodzą na całej swojej powierzchni, zaś grubowarstwowe mogą (ale nie muszą) mieć wycięte laserem przewężenia korygujące ich rezystancję. Z tego powodu ten parametr nie jest podawany przez producentów, można go jedynie próbować oszacować.

Jako przykład takiego oszacowania można przyjąć wykres z **rysunku 2**, który uwidacznia zmianę modułu impedancji rezystora w funkcji częstotliwości. Rezystor 2,2 Ω przejawia cechy typowo indukcyjne, bowiem jego impedancja rośnie wraz ze wzrostem częstotliwości. Z kolei rezystor 10 kΩ zachowuje się bardziej jak pojemność, gdyż jego impedancja maleje. W przypadku rezystora 10 Ω obie te składowe – indukcyjna i pojemnościowa – niemal się znoszą, przez co zachowuje on swoją impedancję niemal równą rezystancji nominalnej w szerokim zakresie częstotliwości.



Rysunek 2. Przebieg modułu impedancji rezystorów w funkcji częstotliwości [2]

użyć jednego elementu w obudowie 0805, a potem się na pewno coś znajdzie w hurtowni... Niestety, rezystory o wartości wyższej niż 10 MΩ bywają trudnodostępne. Ta sama uwaga dotyczy niskich wartości rezystancji. O ile 1 Ω występuje w obudowie 0805 czy 0603, o tyle w 0201 już nie. Przytoczona tabela dotyczy rezystorów grubowarstwowych, więc w razie problemów trzeba sięgnąć po, chociażby, rezystory cienkowarstwowe.

STANDARD ELECTRICAL SPECIFICATIONS									
MODEL	CASE SIZE INCH	CASE SIZE METRIC	POWER RATING P _{10°C} W	LIMITING ELEMENT VOLTAGE MAX. V	TEMPERATURE COEFFICIENT ppm/K	TOLERANCE %	RESISTANCE RANGE Ω	E-SERIES	
CRCW0201...BC	0601	RR 0603M	0.05	30	±200	±0.5	10R to 10M	E96	E96
					-200/+400	±1	1R0 to 9R76		
					±100	±1	4R7 to 1M		
					-200/+400	±5	1R0 to 10M		
CRCW0402...BC	0402	RR 1005M	0.063	50	±100	±1	1R0 to 10M	E96	E96
					-200/+400	±1	1R0 to 9R76		
					±100	±1	1R0 to 10M		
					-200/+400	±5	1R0 to 10M		
CRCW0603...BC	0603	RR 1608M	0.10	75	±100	±1	1R0 to 10M	E96	E96
					-200/+400	±1	1R0 to 9R76		
					±100	±1	1R0 to 10M		
					-200/+400	±5	1R0 to 10M		
CRCW0805...BC	0805	RR 2012M	0.125	100	±100	±1	1R0 to 10M	E96	E96
					-200/+400	±1	1R0 to 9R76		
					±100	±1	1R0 to 10M		
					-200/+400	±5	1R0 to 10M		
CRCW1206...BC	1206	RR 3216M	0.25	200	±100	±1	1R0 to 10M	E96	E96
					-200/+400	±1	1R0 to 9R76		
					±100	±1	1R0 to 10M		
					-200/+400	±5	1R0 to 10M		
CRCW1210...BC	1210	RR 3225M	0.50	200	±100	±1	1R0 to 10M	E96	E96
					-200/+400	±1	1R0 to 9R76		
					±100	±1	1R0 to 10M		
					-200/+400	±5	1R0 to 10M		
CRCW2010...BC	2010	RR 5020M	0.75	400	±100	±1	1R0 to 10M	E96	E96
					-200/+400	±1	1R0 to 9R76		
					±100	±1	1R0 to 10M		
					-200/+400	±5	1R0 to 10M		
CRCW2512...BC	2512	RR 6325M	1.0	500	±100	±1	1R0 to 10M	E96	E96
					-200/+400	±1	1R0 to 9R76		
					±100	±1	1R0 to 10M		
					-200/+400	±5	1R0 to 10M		

Rysunek 3. Wartości rezystorów SMD produkowanych przez firmę Vishay [3]

Dostępny zakres rezystancji

Niby taka niepozorna rzecz, a potrafi napsuć krwi. Nie ma to związku ze schematem z rysunku 1, ograniczenia te wynikają po prostu z technologii, jaką dysponuje dany producent. W notach katalogowych rezystorów są dostępne tabele informujące o produkowanych rezystancjach dla danej wielkości obudowy, choćby jak na **rysunku 3**. Warto też zwrócić uwagę na tolerancję, gdyż ona również może wpływać na rezystancje produkowanych elementów. Nawiasem mówiąc, miło ze strony producenta (tutaj akurat Vishay), że podaje rzeczywistą, spodziewaną rezystancję rezystora „0 Ω” oraz jego wytrzymałość prądową.

Piszę o tym zagadnieniu nieprzypadkowo. Miałem w swojej karierze przypadek, kiedy to do ustalenia częstotliwości pracy generatora potrzebowałem rezystora o rezystancji 15 MΩ. Ponieważ miałem bardzo mało miejsca uznałem, że na etapie projektowania

Item Type	Power Rating at 70°C	Operating Temp. Range	Max. Operating Voltage	Max. Overload Voltage	Resistance Range					TCR (PPM/°C)
					±0.1%	±0.05%	±0.1%	±0.25%	±0.5%	±1%
AR02 (0402)	1/16W	-55 ~ +155°C	50V	100V	49.9Ω - 4.99KΩ					±1, ±2, ±3
					49.9Ω - 20KΩ					±5
					49.9Ω - 100KΩ					±10, ±15
AR03 (0603)	1/16W	-55 ~ +155°C	50V	100V	24.9Ω - 15KΩ					±1, ±2, ±3
					24.9Ω - 60KΩ					±5
					24.9Ω - 100KΩ					±10, ±15
AR05 (0805)	1/10W	-55 ~ +155°C	100V	200V	24.9Ω - 30KΩ					±1, ±2, ±3
					24.9Ω - 100KΩ					±5
					24.9Ω - 1MΩ					±10, ±15
AR08 (1206)	1/8W	-55 ~ +155°C	150V	300V	24.9Ω - 49.9KΩ					±1, ±2, ±3
					24.9Ω - 300KΩ					±5
					24.9Ω - 1.5MΩ					±10, ±15
AR13 (1210)	1/4W	-55 ~ +155°C	150V	300V	24.9Ω - 49.9KΩ					±1, ±2, ±3
					24.9Ω - 300KΩ					±5
					24.9Ω - 1MΩ					±10, ±15
AR10 (2010)	1/4W	-55 ~ +155°C	150V	300V	24.9Ω - 100KΩ					±1, ±2, ±3
					24.9Ω - 300KΩ					±5
					24.9Ω - 1MΩ					±10, ±15
AR12 (2512)	1/2W	-55 ~ +155°C	150V	300V	24.9Ω - 100KΩ					±1, ±2, ±3
					24.9Ω - 300KΩ					±5
					24.9Ω - 1MΩ					±10, ±15

Rysunek 4. Wartości precyzyjnych rezystorów od firmy Viking [4]

Inny producent – Viking – dla swoich cienkowarstwowych rezystorów precyzyjnych podaje jeszcze inne przedziały dostępnych rezystancji (**rysunek 4**). Są one ściśle powiązane z tolerancją. Te najdokładniejsze elementy można kupić w zakresie od 24,9 Ω , zaś maksymalna wartość tego parametru wynosi niecałe pół megoma. Warto mieć to na uwadze podczas projektowania chociażby mikroamperomierzy czy innych układów wymagających takich rezystorów.

Szumy

To już w ogóle temat rzeka. Każdy rezystor wprowadza szum termiczny, którego natura jest dobrze znana i opisana matematycznie. Nic z tym faktem nie zrobimy, możemy co najwyżej chłodzić układy mające szumieć najmniej, co zresztą czasem się robi – wzmacniacze sygnałów pochodzących z np. radioteleskopów pracują chłodzone ciekłymi gazami, aby zmniejszyć wpływ szumu termicznego.

Jednak każdy element rzeczywisty wprowadza też pewne szumy nadmiarowe, które wynikają z jego budowy. Ich konkretna wartość jest trudna do uchwycenia, bowiem zależą one od wielu czynników, głównie od sposobu wykonania ścieżki oporowej i defektów w niej zawartych. To na tychże defektach „rozbijają się” nośniki ładunku elektrycznego, co zaburza ich przepływ, tworząc w efekcie szum dodany do przenoszonego sygnału.

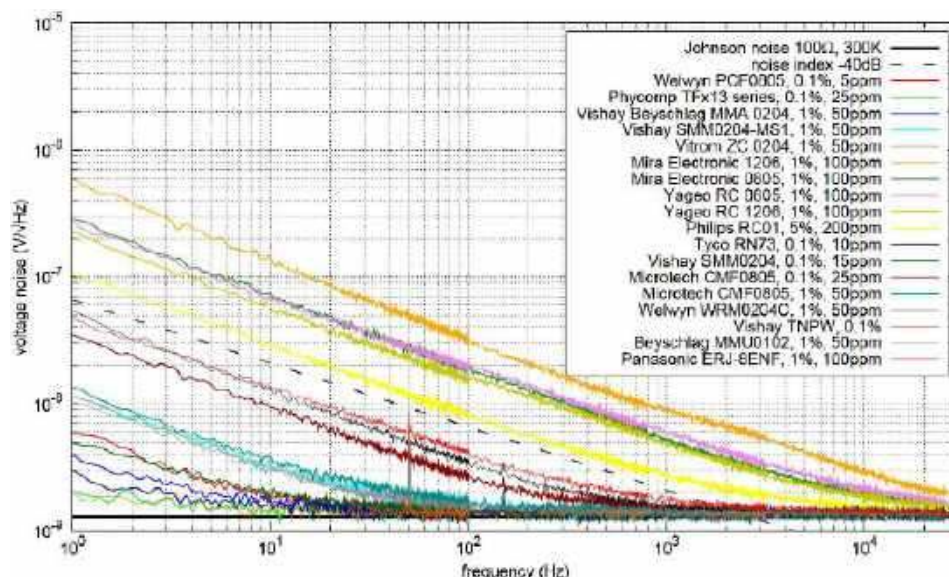
To zjawisko nosi nazwę szumu prądowego, bowiem zależne jest od natężenia prądu płynącego przez rezystor. Na **rysunku 5** znajduje się taki wykres, który pokazuje, jak „szumią” różne rezystory SMD o rezystancji 100 Ω . Można z tego wykresu wywnioskować, iż mniejsze szumy wprowadzają rezystory o niskiej tolerancji, rzędu 0,1%, zwłaszcza w zakresie małych częstotliwości. W miarę wzrostu częstotliwości różnice te stopniowo się zacierają, co oznacza, że do układów niskoszumnych warto stosować rezystory o wąskiej tolerancji wykonania.

Kondensatory

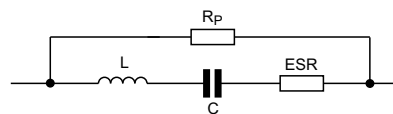
Kondensator, jaki jest, każdy widzi – ot, wiaderko na prąd, jak mawiają niekiedy wykładowcy na początku semestru. Magazynuje ładunek elektryczny, po czym go oddaje. Wszystko jest proste, dopóki działamy w zakresie małych prądów, niskich napięć, niewielkich częstotliwości i w ogóle bez większych wymagań. Uproszczony schemat zastępczy kondensatora znajduje się na **rysunku 6** i mogę z całą stanowczością stwierdzić, że nie wyczerpuje on pełni zjawisk, jakie dotyczą tej grupy podzespołów. Omówię tutaj wyłącznie kondensatory z suchym dielektrykiem, ponieważ elektrolityczne są stosunkowo dobrze opisane w notach katalogowych ich producentów.

Indukcyjność

Znowu...? Niestety. Wszystko, co przewodzi prąd, wytwarza wokół siebie pole magnetyczne, które stanowi magazyn energii. A to już jest podręcznikowy opis indukcyjności. Jak wygląda to zjawisko w przypadku kondensatorów? Na **rysunku 7** znajduje się wykres modułu impedancji kondensatorów MLCC o pojemności 100 nF w różnych obudowach. Taki obrazek jest znany każdemu, kto miał styczność z zagadnieniem określanym jako Power Integrity lub pokrewnymi dziedzinami elektroniki. Powyżej częstotliwości rezonansowej, kondensator przestaje być pojemnością, a staje się indukcyjnością, więc jego zdolności filtracyjne są wówczas mocno wątpliwe. Dlatego stosuje się wiele kondensatorów o różnych pojemnościach,



Rysunek 5. Wykres napięcia szumu prądowego różnych rezystorów SMD o wartości 100 Ω [5]



Rysunek 6. Prosty schemat zastępczy kondensatora [1]

połączonych równolegle. Wtedy ich częstotliwości rezonansowe nie nakładają się, więc każdy z nich działa poprawnie w innym zakresie częstotliwości. Potwierdza to wykres z **rysunku 8**, na którym widać przebieg modułu impedancji kondensatorów w takich samych obudowach 1206, lecz o różnych pojemnościach. Na podstawie wykresu można też wywnioskować, że im wyższa jest pojemność kondensatora, tym niższa jego częstotliwość rezonansu własnego.

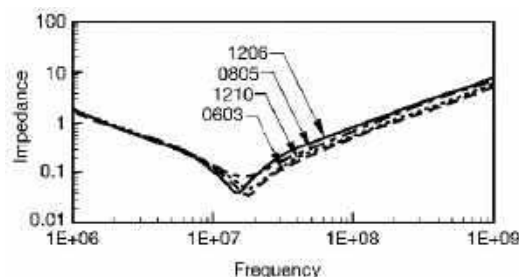
W przypadku pomiarów, których wyniki zebrano na rysunku 7 można zauważyć, że częstotliwości rezonansowe tych kondensatorów nieznacznie się różnią. Obliczenie ich teoretycznych indukcyjności daje następujące wartości [6]:

- 0603: 870 pH;
- 0805: 1050 pH;
- 1206: 1200 pH;
- 1210: 980 pH.

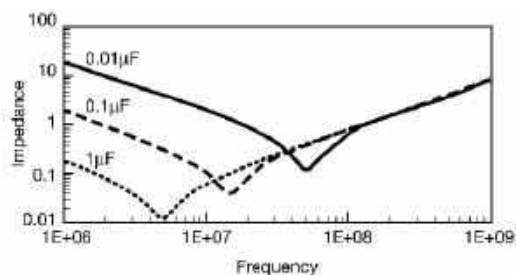
Na tej podstawie nie można, niestety, stwierdzić jasnej zależności między indukcyjnością kondensatora a rozmiarem jego obudowy – element w obudowie 1210 wyłamuje się z trendu, zapewne przez to, że jego długość i szerokość są najbardziej zbliżone do siebie spośród wszystkich czterech typów obudów (stosunek 12:10).

Efekt ferroelektryczny

Jak to się dzieje, że taki ledwie dostrzegalny na płytce kondensator ma pojemność nawet kilkunastu mikrofaradów? Fizyka mówi, że trzeba zwiększyć powierzchnię okładek lub... zwiększyć przenikalność elektryczną dielektryka. I tym właśnie operują producenci,



Rysunek 7. Przebieg modułu impedancji kondensatorów MLCC 100 nF w funkcji częstotliwości [6]



Rysunek 8. Przebieg modułu impedancji kondensatorów w obudowach 1206, o różnych pojemnościach, w funkcji częstotliwości [6]

tworząc kondensatory ceramiczne o pojemnościach porównywalnych z kondensatorami elektrolitycznymi i do tego o znacznie mniejszych gabarytach. Gdzie więc tkwi haczyk i po co stosować duże elektrolity? Odpowiedzi może być wiele, ale jedna z nich znajduje się na wykresie z **rysunku 9**.

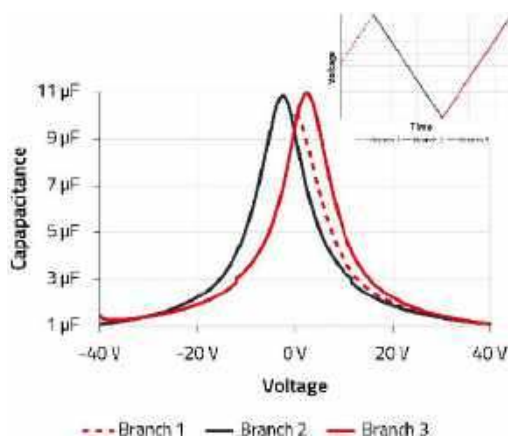
Do kondensatora MLCC o nominalnej pojemności 10 μF zostało przyłożone stałe napięcie o przebiegu trójkątnym, jak widać w prawym górnym rogu rysunku 9. Na pierwszym zboczku, narastającym od zera, pojemność maleje w miarę wzrostu napięcia, nawet do 1 μF – czyli do zaledwie... 10% wartości początkowej! Kiedy napięcie opada – krzywa w kolorze czarnym – pojemność znowu wzrasta, lecz już z nieco innym przebiegiem niż poprzednio. Po silnym spolaryzowaniu w przeciwną stronę, drugie zbocze narastające składowej stałej napięcia wywołuje zmianę pojemności o jeszcze innym przebiegu.

Takie zachowanie kondensatora jest tłumaczone ustawianiem się domen elektrycznych w materiale dielektryka zgodnie z wektorem pola elektrycznego między okładkami. Im silniejsza polaryzacja napięciem stałym, tym więcej z nich ulegnie obróceniu, co wpłynie negatywnie na pojemność kondensatora, tj. zmniejszy ją. Jednak przywrócenie z powrotem tej samej wartości napięcia polaryzującego kondensator, co poprzednio, nie przywróci jego poprzedniej pojemności, gdyż część domen utrzyma swoje położenie. Jest to efekt analogiczny do tego, który spotykamy w ferromagnetykach.

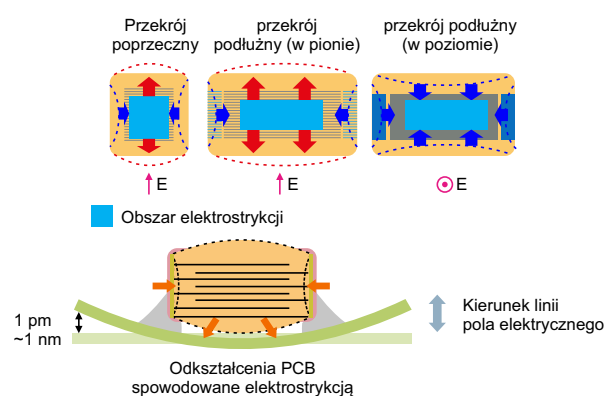
Jakie wnioski możemy z tego wyciągnąć? Kondensatorów MLCC nie stosuje się tam, gdzie zależy nam na pojemności stabilnej w funkcji napięcia, czyli na przykład w generatorach RC. Z kolei przy separacji składowej stałej (na przykład pomiędzy stopniami wzmacniającymi) trzeba pamiętać o spadku ich pojemności po spolaryzowaniu lub sięgnąć po inny kondensator i nie zawracać sobie głowy dodatkowymi problemami. Takie podejście uzasadnia również następny punkt, czyli...

Mikrofonowanie

Do tej właściwości kondensatorów MLCC producenci nie chcą się przyznawać, lecz ona istnieje. Wpadłem na jej trop wiele lat temu, kiedy zastosowałem kondensator 100 nF w obudowie 1206 (na odpowiednie napięcie, chyba 200 V) do sprzęgania



Rysunek 9. Zależność między pojemnością kondensatora MLCC a polaryzującą go składową stałą napięcia [7]



Rysunek 10. Efekt piezoelektryczny w kondensatorach MLCC [8]

stopni wzmacniacza lampowego... Chciałem być taki nowoczesny, bo esemde i w ogóle miniaturyzacja, a tymczasem wyszedł dramat.

W dokumencie [8] została poruszona kwestia wydawania dźwięków przez kondensatory MLCC w obwodach, które pracują impulsowo, na przykład w przetwornicach. Efekt piezoelektryczny powoduje odkształcanie laminatu, jak widać na **rysunku 10**. Pamiętajmy jednak, że efekt ten jest odwracalny, więc skoro kondensator odkształca się po przyłożeniu do niego napięcia, to co dzieje się z nim, kiedy jego obudowa jest poddawana naprężeniom? Generuje ładunek elektryczny, czyli zmienia napięcie na swoich okładkach.

O ile w odsprężaniu zasilania w ogóle nie interesuje nas ten problem, to przy sprzęganiu ze sobą stopni układu audio taki kondensator potrafi być źródłem buczenia o częstotliwości 50 Hz, które po powierzchni laminatu przenosi się z transformatora zasilającego. Żadne ekrany elektryczne czy magnetyczne na to nie pomogą, żadna filtracja zasilania też nie – taki element trzeba po prostu z danego miejsca układu usunąć i zastąpić innym.

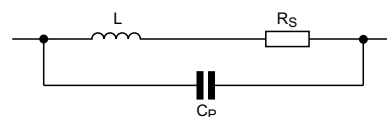
Cewki

Na sam koniec zostawiłem te najdziwniejsze, chyba najmniej lubiane podzespoły biernie. Niby działają analogicznie jak kondensator, ale są przy tym tak bardzo niedoskonałe, że niekiedy szkoda gadać. Najbardziej ogólny schemat zastępczy cewki znajduje się na **rysunku 11**. Oprócz właściwej indukcyjności L , mamy też rezystancję szeregową R_S i pojemność wewnętrzną C_P . Ich pochodzenie jest dosyć jasne: rezystancja bierze się z ograniczonej przewodności drutu lub innego materiału, którym została nawinięta dana cewka, zaś pojemność wynika z pojemności pomiędzy wyprowadzeniami obudowy, jak i między poszczególnymi zwojami.

Rezonans własny

Skoro mamy indukcyjność i pojemność, to mamy również rezonans elektryczny między nimi. Dobroć cewki jest ograniczona przez jej pasożytniczą rezystancję, więc przebieg impedancji w funkcji częstotliwości nie ma silnie zaokrąglonego wierzchołka. Mimo to ów rezonans wciąż istnieje i powyżej punktu granicznego cewka zachowuje się bardziej jak kondensator. Kiedy wstawimy do układu filtr LC, to zamiast filtrować różne „śmieci” z obwodów zasilania, ten zaczyna je przepuszczać jeszcze łatwiej!

Warto zauważyć, że częstotliwość rezonansu własnego nie musi być jakaś kosmicznie wysoka, co udowadnia lista znajdująca się na **rysunku 12**. Wymienione na niej dławiki są produkowane w obudowach 1206, zaś czerwoną pętlą zaznaczyłem te wartości, które najczęściej wstawia się w tor zasilania. W wielu



Rysunek 11. Prosty schemat zastępczy cewki [1]

poradnikach dla początkujących można znaleźć filtry LC powstawiane bezmyślnie w tor zasilania mikrokontrolera, który jest taktowany zegarem o częstotliwości kilkunastu megaherców. Nie trzeba być orłem i sokołem, by się przekonać, że na tej częstotliwości, bądź przynajmniej przy drugiej czy trzeciej harmonicznej, taka cewka o indukcyjności kilku mikrohenrów jest bezużyteczna.

Zostawmy zatem te biedne cewki w tych miejscach, gdzie spisują się zdecydowanie lepiej – przede wszystkim w układach dopasowujących impedancję. Do filtracji zasilania zdecydowanie lepiej nadają się koraliki SMD, które wprowadzają straty w składowej zmiennej prądu, a nawet proste filtry RC (w przypadku obwodów pobierających prąd o niskim natężeniu).

Prąd pracy

To zagadnienie dotyczy głównie dławików mocy, które pracują ze składową stałą prądu o relatywnie wysokim natężeniu. O co chodzi? Jeszcze kilka...kilkanaście lat temu dosyć powszechne (przynajmniej moim zdaniem) było zjawisko podawania przez producentów indukcyjności przy zerowym lub bardzo niskim natężeniu prądu stałego płynącego przez taką cewkę. Ponieważ rdzenie dławików mocy są wykonane z ferromagnetyków, toteż ich indukcyjność była wówczas atrakcyjnie wysoka.

Z kolei podawany szumnie tuż obok maksymalny prąd pracy dotyczył maksymalnego prądu, jaki mógł płynąć przez drut uzwojenia, bez ryzyka przegrzania go wskutek wydzielanego w nim ciepła. Tyle, że wskutek podmagnesowania rdzenia, z tej indukcyjności niewiele już zostawało, a w skrajnych przypadkach ów parametr był wyższy od prądu nasycenia, kiedy to indukcyjność spada – teoretycznie – do zera.

Jaką zatem wartość prądu przyjąć? Nowe karty katalogowe są już, na całe szczęście, pisane w inny sposób niż kiedyś, mniej marketingowy. Wyraźnie zaznacza się w nich maksymalny prąd wynikający z wytrzymałości materiału uzwojenia, zaś na pierwszym miejscu z reguły podaje się maksymalny prąd pracy, czyli taki, przy którym indukcyjność dławika pozostaje zachowana. Niekiedy też, jak na **rysunku 13**, podawany jest elegancką definicję pojęcia „prąd nasycenia”. W tym wypadku (nota katalogowa dławików marki Kyocera) jest to prąd, przy którym indukcyjność spada o 10% względem nominalnej, co w zdecydowanej większości przetwornic i innych układów dużej mocy nie grozi ich uszkodzeniem bądź nieprawidłową pracą. Pociągający jest fakt, że obecnie wielu producentów przyjmuje taki sam sposób definicji prądu nasycenia, co ułatwia porównywanie elementów.

Podsumowanie

Elementy biernie, pomimo swojej pozornej prostoty, rządzą się niekiedy bardzo złożonymi prawami. Do tego stopnia, że rezystory potrafią zachowywać się jak cewki, kondensatory tracić swoją pojemność, a cewki – zaburzać pracę całego układu, którego funkcjonowanie miały poprawiać. Dlatego na rynku nie mamy tylko jednego rodzaju elementu, zaś całą ich gamę, w różnych wykonaniach i obudowach.

Tym artykułem chciałem przybliżyć niektóre zagadnienia związane z podzespołami biernymi, które rzadko są omawiane na różnego rodzaju kursach. Ich nieznanomość nie przeszkadza, kiedy próbuje się uruchomić multiwibrator astabilny na dwóch tranzystorach, lecz przy projektowaniu wzmacniacza mikrofonowego lub

STANDARD ELECTRICAL SPECIFICATIONS								
PART NUMBER	INDUCTANCE (μH)	TOL. (%)	TEST FREQ. (MHz)		Q MIN.	SRF MIN. (MHz)	DCR MAX. (Ω)	RATED DC CURRENT (mA)
			L AND Q					
ILSB1206ER47NM	0.047	20	50	20	368	0.16	300	
ILSB1206ER62NM	0.068	20	50	20	322	0.25	300	
ILSB1206ER10K	0.10	10	25	20	271	0.25	250	
ILSB1206ER12K	0.12	10	25	20	253	0.30	250	
ILSB1206ER15K	0.15	10	25	20	230	0.30	250	
ILSB1206ER18K	0.18	10	25	20	213	0.40	250	
ILSB1206ER22K	0.22	10	25	20	196	0.40	250	
ILSB1206ER27K	0.27	10	25	20	173	0.50	250	
ILSB1206ER33K	0.33	10	25	20	157	0.60	250	
ILSB1206ER39K	0.39	10	25	20	136	0.50	200	
ILSB1206ER47K	0.47	10	25	20	144	0.60	200	
ILSB1206ER56K	0.68	10	25	20	121	0.80	150	
ILSB1206ER100K	1.0	10	10	45	87	0.40	100	
ILSB1206ER12K	1.2	10	10	45	75	0.50	100	
ILSB1206ER15K	1.5	10	10	45	69	0.50	50	
ILSB1206ER18K	1.8	10	10	45	64	0.50	50	
ILSB1206ER22K	2.2	10	10	45	58	0.50	50	
ILSB1206ER33K	3.3	10	10	45	46	0.70	50	
ILSB1206ER39K	3.9	10	10	45	44	0.80	50	
ILSB1206ER47K	4.7	10	10	45	41	0.90	50	
ILSB1206ER56K	5.6	10	4	45	37	0.70	25	
ILSB1206ER68K	6.8	10	4	45	34	0.80	25	
ILSB1206ER100K	8.2	10	4	45	30	0.90	25	
ILSB1206ER100K	10	10	2	45	26	1.00	25	
ILSB1206ER120K	12	10	2	45	26	1.05	15	
ILSB1206ER150K	15	10	1	45	22	0.70	5	
ILSB1206ER180K	18	10	1	45	21	0.70	5	
ILSB1206ER220K	22	10	1	35	19	0.90	5	
ILSB1206ER270K	27	10	1	35	17	0.90	5	
ILSB1206ER330K	33	10	1	35	15	1.05	5	

Rysunek 12. Zestawienie indukcyjności cewek z ich częstotliwościami rezonansu własnego [9]

Codes	L (μH)	Tolerance	Test Condition	DCR (Ω) max.	I sat. (A) max*
4R7	4.7	M	100KHz, 0.1V	0.056	4.20
6R8	6.8	M	100KHz, 0.1V	0.060	3.00
100	10	M	100KHz, 0.1V	0.065	2.70
150	15	M	100KHz, 0.1V	0.072	2.30
220	22	M	100KHz, 0.1V	0.074	1.80
330	33	M	100KHz, 0.1V	0.085	1.60
470	47	M	100KHz, 0.1V	0.092	1.30
680	68	M	100KHz, 0.1V	0.094	1.10
100	100	M	100KHz, 0.1V	0.099	0.87
150	150	M	100KHz, 0.1V	0.094	0.74
220	220	M	100KHz, 0.1V	1.00	0.56
330	330	M	100KHz, 0.1V	2.15	0.60
470	470	M	100KHz, 0.1V	3.30	0.40
680	680	M	100KHz, 0.1V	4.40	0.35
100	1000	M	100KHz, 0.1V	7.00	0.25

*Saturation Current: The current when the inductance becomes 10% lower than the nominal value. (for 25°C)

Rysunek 13. Przykład jasnej definicji prądu nasycenia dławika mocy [10]

innego „klocka” analogowego potrafią doprowadzić do wybuchu furii w pracowni elektronicznej. Zła wiadomość: ci, którzy parają się wyłącznie cyfrówką i myślą, że zagadnienia podejrzanych zachowań elementów biernych ich nie dotyczą, są w błędzie. Fajne i tanie kondensatory MLCC kiepsko radzą sobie z filtrowaniem zasilania, zaś niefortunnie użyty dławik potrafi zepsuć zasilanie całego układu cyfrowego.

Nie ma jednak tego złego, co by na dobrze nie wyszło: do projektowania trzeba podchodzić po prostu świadomie, mając na uwadze zarówno wady, jak i zalety poszczególnych elementów. Kondensatory MLCC czy miniaturowe dławiki są, jakie są, ale mogą mieć sporo zastosowań we współczesnych układach elektronicznych, byleby stosować je tam, gdzie naprawdę jest ich miejsce.

Michał Kurzela, EP

Źródła:

- [1] <https://www.ti.com.cn/cn/lit/an/sloa027/sloa027.pdf>
- [2] https://www.koaglobal.com/product/library/resistor/characteristics?sc_lang=en
- [3] https://www.tme.eu/Document/622e28308c06d160637f2b14751ba16b/Data%20Sheet%20CRCW_BCe3.pdf
- [4] <https://www.tme.eu/Document/bf44a96e1b73416f868b2cf96e245ffd/AR%20REV%20F5%2020240522.pdf>
- [5] https://dcc.ligo.org/public/0002/T0900200/001/current_noise.pdf
- [6] <https://kyocera-avx.com/docs/techinfo/CeramicCapacitors/parasitic.pdf>
- [7] https://www.onsemi.com/components/media/0753710v410%20ANP114a_Polarization%20DC%20Bias%20MLCC_EN.pdf
- [8] <https://www.ti.com/lit/ta/ssztb09/ssztb09.pdf>
- [9] <https://www.vishay.com/docs/34029/ilsb1206.pdf>
- [10] <https://catalogs.kyocera-avx.com/SMD-PowerInductors.pdf>



FN-SWM10

Zgrzewarka do ogniw – spawarka punktowa z kolorowym wyświetlaczem i funkcją powerbank FNIRSI SWM10



FN-DPOS-350P

Dwukanałowy oscyloskop 350 MHz, FNIRSI DPOS350P



FN-2C53T

Dwukanałowy oscyloskop z multimetrem i generatorem 50 MHz FNIRSI 2C53T

BESTSELLERY sklepu AVT – sklep.avt.pl

Mierniki Testery FNIRSI

Rabat dla Czytelników EP przy zakupie podaj kod **EP2505FN**

-3%

Rabat dla Prenumeratorów EP przy zakupie podaj numer prenumeraty

-6%



FN-LCR-ST1

Miernik pęsetowy, tester elementów FNIRSI LCR-ST1



FN-LCR-P1

Tester elementów FNIRSI LCR-P1



FN-HRM10

Tester rezystancji wewnętrznej akumulatorów FNIRSI HRM-10



FN-G1200

Mikroskop cyfrowy G1200 z wyświetlaczem 7 cali, powiększenie ×1200, tryb foto/video



FN-DWS200-F245

Stacja lutownicza 200 W z kolbą F245, FNIRSI DWS200



FN-1014D

Oscyloskop dwukanałowy 100 MHz; Generator sygnału DDS, FNIRSI 1014D

Poniższym artykułem rozpoczynamy na łamach EP nowy cykl felietonów, w których oddamy głos nie tylko naszym stałym Autorom, ale także Czytelnikom obdarzonym lekkim piórem i zainteresowanym podzieleniem się swoimi przemyśleniami na temat szeroko pojętej technologii. Zapraszamy do współtworzenia nowej rubryki!



Retromania, czyli jak zarobić na starociach

Moda na stare przedmioty trwa w najlepsze od lat. O ile jeszcze można zrozumieć zamięłowanie – zarówno audiofilów, jak i gitarzystów – do charakterystycznego brzmienia lampowego wzmacniacza, to już powrót do łask innych technologii i urządzeń z czasów „słusznie minionych” jest dla mnie zjawiskiem zadziwiającym. Dlaczego ktokolwiek chciałby kupić i używać sprzętu oraz nośników analogowych pod każdym względem gorszych od współczesnych rozwiązań? Ba, jakimś cudem niektórzy są przekonani, iż stara taśma magnetyczna czy płyta z tworzywa brzmi lepiej, niż cyfrowy plik. I na to brzmienie są gotowi wydać absurdalne kwoty.

Stary, (nie-)dobry nośnik analogowy

Płyty winylowe mają swój szczególny urok fizycznego nośnika, na którym wręcz widać nagranie. Technologia jest prosta do zrozumienia: ot, wytłoczony rowek o zmiennej głębokości i szerokości. W tych zmianach ukryty jest dyskretny sygnał analogowy, który wystarczy odczytać igłą sprzężoną z przetwornikiem

piezoelektrycznym lub magnetycznym, wzmacnić, dopasować pasmo do krzywej RIAA i ponownie wzmacnić. Oczywiście płyta musi się obracać ze stałą prędkością $33\frac{1}{3}$ lub 45 obrotów na minutę. Od dekad dostępne są dedykowane układy stabilizatorów obrotów oparte na rezonatorach kwarcowych. Wcześniej realizowano to zadanie projektując odpowiednie przekładnie i silniki synchroniczne prądu zmiennego, których obroty zależały od częstotliwości napięcia sieciowego. Skoro technologia jest tak prosta, to dlaczego nowy gramofon polskiej produkcji kosztuje dwanaście tysięcy złotych? Tymczasem gramofon znanej marki Yamaha kosztuje maksymalnie $\frac{1}{4}$ tej ceny, zaś sprzęt firmy Audio-Technica jest dostępny w cenie około 7...10% ceny audiofilskiego „Fryderyka”. Czym zatem różni się gramofon za niespełną tysiąc złotych od tego za dwanaście tysięcy? Ceną. Serio, urządzenie marki Audio-Technica może być nieco gorszej jakości, z większą ilością plastiku i bez drewnianej bazy, ale funkcjonalnie niczym się nie różni od gramofonu rodzimej produkcji kosztującego dziesięć razy więcej. Polski producent będzie się zarzekał, że to nieprawda, że ich napęd, projekt ramienia z igłą i elektronika są znacznie lepsze, ale w ostatecznym rozrachunku te elementy nie robią różnicy – sam nośnik jest bowiem na tyle niedoskonały, że typowe zabiegi „audiofilskie” nie poprawiają jego jakości. Nie pomoże wzmacniacz z SNR 120 dB, gdy zakres dynamiki

plyty to maksymalnie 70 dB (nowa płyta, nowa, dobrej jakości igła magnetyczna, doregulowana siła nacisku). Nie pomoże idealnie płaski i wyważony talerz, gdy same płyty nigdy płaskie i wyważone nie są – proces produkcji w końcu polega na brutalnym miażdżeniu tworzywa w rozgrzanej prasie, co wytłacza wzór z matrycy w tworzywie, a potem na wystudzeniu tegoż, co nigdy nie zachodzi równomiernie i tworzy mikroskopijne naprężenia i odkształcenia materiału. Każdy specjalista od automatyki przemysłowej czy systemów kontroli potwierdzi, iż stworzenie napędu trzymającego stałe obroty z użyciem sprzężenia zwrotnego to nie jest duży problem i potrafi to chiński serwowator za kilkadziesiąt złotych, przeznaczony do napędzania dużych cięśliwych od płyty winylowej maszyn – z precyzją często liczoną w ułamkach milimetra. Budowa specjalnego napędu do prostego mechanicznie gramofonu mija się zatem z celem i nie usprawiedliwia wysokiej ceny gotowego urządzenia. Ktoś jednak kupuje te urządzenia. I nie, „Fryderyk” nie jest jedynym gramofonem za grube tysiące – wiele małych, niszowych firm produkuje sprzęt „audiofili” do odtwarzania prasowanych krążków z tworzywa. Moda na winyle trwa w najlepsze od czasów pierwszych komercyjnych magnetofonów szpulowych. I choć przez pewien czas płyty winylowe stały się na tyle niszowym produktem, że wiele tłocznii przestało działać, to od dekady nowe, często małe manufaktury odkupują stare maszyny, przywracają je do życia i wznawiają tłoczenie płyt, często w małych seriach. Proces jest na tyle skomplikowany, że każda maszyna tłocząca ma swoje „dziwactwa” i nawet dwie identyczne prasy mogą wymagać innych ustawień siły nacisku, temperatury i czasu tłoczenia. W końcu tworzone jest fizyczne medium. I choć cena produkcji płyt winylowych jest nadal wysoka, a urządzenia odtwarzające bywają drogie, nikogo to nie zraża i płyta winylowa przeżywa swój renesans.

Odbija to się też na rynku wtórnym – stare gramofony, produkowane przez dawne przedsiębiorstwo Unitra Fonica, osiągają ceny od tysiąca do ponad czterech tysięcy złotych. Gramofony marek zagranicznych, jak na przykład Technics, potrafią kosztować tyle, co nowy „Fryderyk”, mimo że mają więcej lat, niż ja. Nawet fatalne, niszczące płyty „Bambino” potrafi kosztować dwieście złotych, gdy za mniej niż 500 można nabyć chiński gramofon, który będzie o rząd wielkości lepszy. Najlepszą rzeczą, jaką można zrobić z „bambino” jest prosty „piecyk gitarowy” o miłym dla ucha, lampowym brzmieniu. Odtwarzanie na nim płyt tylko je niszczy.

No to może wspomniane wcześniej taśmy magnetyczne? Do wyboru jest wiele formatów, ale dwa główne to taśmy szpulowe i kompaktowe kasety magnetofonowe. Te pierwsze były niegdyś podstawowymi nośnikami stosowanymi w studiach nagraniowych i rozgłośniach radiowych oraz wszędzie indziej, gdzie potrzebne były: duży zakres dynamiki i wysoka jakość dźwięku. Kasety magnetofonowe miały gorsze parametry, za to były znacznie wygodniejsze i pozwalały na odtwarzanie za pomocą przenośnych urządzeń mieszczących się w kieszeni albo przypiętych klipssem do paska. Kaseta pozwoliła na montaż odtwarzaczy w samochodach (rzecz wcześniej raczej niespotykana). Istniały gramofony samochodowe, ale były to produkty duże, ciężkie, skomplikowane i niewygodne, wymagające fizycznego przykręcenia płyty do talerza, by ta się nie chwiała i by igła, zamocowana na sztywnym ramieniu, zachowała idealną odległość od nośnika. Kaseta magnetofonowa miała też sporą pojemność: 60 lub 90 minut (były też kasety 120-minutowe, ale ich taśmy były delikatne i łatwo się rwały). Na początku pojawiły się magnetofony stricte użytkowe – reporterskie i dyktafony. Dość szybko jednak format kasety kompaktowej trafił w ręce miłośników muzyki, gdyż oferował faktycznie lepsze parametry od winyli, w wygodniejszej formie. Dość szybko rynki światowe zalały magnetofony różnej klasy, od modeli Hi-Fi, po towary podlegszego gatunku. Co ciekawe, w PRLu na początku produkowano magnetofony mechanicznie dobrze zaprojektowane, gdzie mechanizm był wykonany z metalu, podobnie jak koła zębate przekładni. Te mechanizmy licencyjne

były niezniszczalne. Potem w ramach racjonalizacji, i to nie tylko w Polsce Ludowej, ale też w krajach zachodnich, metal zastąpił plastik, koła z mosiądzu zamieniono na poliamidowe, a nawet silniki elektryczne doczekały się integracji ze stabilizatorami obrotów. Teraz magnetofony od Unitry może nie dorównują cenowo gramofonom, ale nadal podłej jakości urządzenie z lat 80. potrafi kosztować kilkaset złotych. W tym jednak wypadku nowe produkty, które mają zaspokoić potrzeby fanów retro, są produkowane nie na bazie modeli Hi-Fi, ale na podstawie mechanizmów z końca ery tajwańskiej produkcji magnetofonów, zredukowanych do 90% tworzywa, często z monofoniczną głowicą słabej jakości. W tej sytuacji jednak trzeba sięgać po sprzęt używany, w tym po wcześniejsze produkty starej Unitry. Modele późniejsze mają problemy natury „dentystycznej” – koła zębate z tworzywa często tracą zęby po 40 latach od daty produkcji.

Mam kolegę, który zbiera taśmy magnetofonowe. Nie wiem, w jakim celu, ale ma ich sporą kolekcję, w tym wyrobów polskiego Stilonu. O ile same opakowania szpul czy kaset mają bardzo ciekawe wzornictwo, o tyle zawartość wypadła gorzej. Nawet w czasach świetności zdarzało się, że taśmy traciły swoją żelazową powłokę, brudząc mechanizm i głowicę. Po wielu latach od daty produkcji nośniki te są w jeszcze gorszym stanie. Taśmy marek zagranicznych, jak na przykład TDK, trzymały się o wiele lepiej. Warto też dodać, że same taśmy miały różne materiały magnetyczne, co wpływało na jakość nagrania i cenę. Taśmy typu I, czyli najpopularniejsze, jako materiały magnetyczne używały tlenku żelaza – były relatywnie tanie i miały wystarczająco dobrą jakość dźwięku, choć gorzej reprodukowały tony wysokie. Taśmy typu II używały dwutlenku chromu albo tlenku żelaza domieszkowanego kobaltom, co poprawiało zakres dynamiki i przenoszenie tonów wysokich, oczywiście kosztem ceny. Taśmy typu III łączyły warstwę tlenku żelaza i warstwę dwutlenku chromu, co dawało zalety taśm typu I i II, ale trudniej było doregulować dla nich prąd podkładu (niewielki sygnał wstępnie magnetyzujący taśmę), przez co nie zyskały dużej popularności. Taśmy typu IV używały czystego metalu: żelaza, chromu lub niklu, oferując najlepsze parametry dźwięku kosztem ceny. Tak czy inaczej zakres dynamiki kaset magnetofonowych wynosił typowo 50...56 dB i wynikał z relatywnie małej szerokości samej taśmy oraz niskiej prędkości posuwu. Do tego natura nośnika magnetycznego generowała dodatkowy szum, który był wyraźnie słyszalny w cichszych partiach nagrania. Ray Dolby opracował na to patent w formie systemu redukcji szumu Dolby A i Dolby B – ten drugi trafił na rynek konsumencki, ale działał tylko wtedy, gdy kaseta była nagrana i odtwarzana z użyciem magnetofonu wspierającego ten system. Dolby A był systemem bardziej skomplikowanym i przez to występował w sprzęcie profesjonalnym. Kolejne systemy: Dolby C, S, SR i X poprawiały jakość nagrań, ale pojawiły się za późno, by pokonać zyskującą popularność płytę CD.

Zagraniczni i krajowi audiofile są gotowi wydać krocie na wspomnianą wcześniej taśmę szpulową. Te zaś mają szereg zalet: zakres dynamiki typowo 80 dB, sporą długość i cztery ścieżki do wyboru. Magnetofon szpulowy marki Revox obecnie może kosztować tyle, co samochód – urządzenia takie występowały głównie w studiach nagraniowych, a na rynku wtórnym działa kilka firm, które skupiają te magnetofony, dogłębnie je remontują i sprzedają zapaleńcom za ciężkie pieniądze. Nie widzę w tym sensu, ale... kto bogatemu zabroni?

Czytelniku, jeśli jesteś miłośnikiem winyli czy kaset, to ostrzegam, iż zamierzam teraz szargać świętości.

Nośniki analogowe są do bani! Zwłaszcza jeżeli porównamy je z cyfrowymi. Płyta CD-Audio ma zakres dynamiki 93 dB, podobnie jak dostępna na rynku konsumenckim od 1987 roku taśma DAT (Digital Audio Tape). Ta druga miała trudne początki, bo wydawcy muzycy bali się bezstratnego kopiowania – kopiowanie nośników analogowych jest zawsze procesem stratnym. Potem na rynku

pojawiały się pliki MP3, a wkrótce po nich odtwarzacze tychże. Część melomanów oraz wszyscy wydawcy zadrżeli, choć z różnych powodów. Wydawcy z powodu piractwa, melomani – z powodu jakości. Ci drudzy zakrzyknęli przeto: „Empetrójki masakrują dźwięk i są do bani” i mieli w tym rację. Format działa przez usunięcie tych informacji, których nie słyszymy, a resztę enkodując i kompresując. Na początku lat dwutysięcznych, gdy Internet był wolny i drogi, a dyski twarde małe, muzyka była enkodowana przez piratów dość agresywnie, często z przepływnością 128 kbps (lub mniejszą) i z redukcją jednego kanału tylko do informacji o różnicach, co jeszcze bardziej pogarszało jakość dźwięku. Teraz cyfrowe formaty stracone i bezstratne są wszędzie, a jakością przewyższając nawet te wczesne nośniki cyfrowe. Audiofil w ślepych teście nie odróżnia jednych od drugich. Dyski twarde, SSD czy karty pamięci są tanie jak barszcz, a wydawcy sprzedają cyfrową muzykę w formie usługi streamingowej, kontrolując rynek i walcząc z piractwem przez oferowanie wygodniejszego rozwiązania. Nośniki analogowe przegrały i tyle. Wrócenie do tych technologii to sztuka dla sztuki i czysty, akustyczny masochizm.

Cmentarzisko zapomnianych technologii

Z zamiłowania jestem historykiem postępu naukowego i technologicznego. Fascynują mnie stare komputery i chciałbym kilka posiadać, jeśli nie w oryginale, to w formie funkcjonalnych replik. I nie mam na myśli maszyn sprzed dwudziestu czy trzydziestu lat, lecz komputery z lat 50., 60. i 70. ubiegłego wieku. Komputery, dla których 128 kilobitów pamięci to było dużo, a dysk o pojemności 20 MB służył wielu użytkownikom. Urzekła mnie magia migających lampek, programowania za pomocą przełączników i możliwość widzenia na własne oczy, jak program wykonuje się za każdym razem, gdy wciśnię się przycisk lub pstryknie przełącznikiem „Single Step”. Współczesny komputer to przy tym niepojęta, czarna skrzynka, w której roją się dziesiątki programów i setki procesów, których nie rozumiem i nie zrozumie. Stare komputery były o wiele prostsze. Co nie zmienia faktu, że nie zamienię mojego peceta, na którym piszę te słowa, na w pełni sprawny Odrę 1305 czy DEC PDP-8. Choć kompletny i sprawny komputer IBM z serii System/360 już by mnie kusił.

Dziwi mnie natomiast trwająca od jakiegoś czasu moda na retrokomputery i retrokonsole. Po co kupować kosztowną reprodukcję konsoli Atari 2600 czy komputera Commodore 64, w których to siedzi procesor ARM albo/i układ FPGA, emulujące oryginał i oferujące garść gier z epoki? Za 200...300 złotych można kupić kieszonkową konsolkę z Chin, z moralnie wątpliwą biblioteką dziesiątek tysięcy gier na każdą konsolę, każdy domowy komputer i każdy automat arcade, od lat 70. po czasy PlayStation 2 czy Nintendo Wii. Zresztą nie trzeba nawet tego robić, bo dostępne są programowe emulatory i biblioteki pirackich ROMów czy archiwa obrazów taśm i dyskietek. Nie trzeba wydawać ani złotówki, by pokazać dzieciom, jak w 1986 roku wyglądała przełomowa grafika w grach (ten koszmar możemy im zaserwować za darmo, by potem zrozumiały, dlaczego ich rodzice nie do końca są normalni – to przez trudne dzieciństwo).

Istnieją jednak ciekawe projekty z pogranicza retro i współczesności, dla przykładu: Commander X16 oraz OtterX. Oba „retrokomputery” używają mikroprocesora W65C02 firmy WDC, będącego współczesnym klonem układu MOS 6502, znanego z komputerów Commodore czy konsol NES (w Polsce zwanych Pegazusami i sprowadzanych we wczesnych latach 90. z Tajwanu przez założycieli marki Hoop) oraz wielu innych komputerów, konsol i urządzeń przemysłowych. Commander X16 wiernie odtwarza doświadczenia używania komputera Commodore 64, ale oferuje przy tym współpracę ze standardowymi klawiaturami i myszkami, gniazda kontrolerów SNES i wiele innych udogodnień. Projekt ma dwie wady: wysoką cenę (349 dolarów + koszty wysyłki) i konieczność użycia oryginalnego układu Yamaha YM2151 (syntezator FM OPM). OtterX to tańszy klon tego projektu, na mniejszej płycie (mini-ITX), w formie zestawu

do samodzielnego montażu (za 275 dolarów + koszty wysyłki). Twórca jednak udostępnił całą dokumentację, więc w teorii każdy może sobie zamówić płytkę u swojego ulubionego producenta w Chinach, a następnie samemu złożyć całość. Autor poprawił kilka błędów oryginalnego projektu, jak i kartę VERA, oznaczoną jako VERAX. Twórca projektu oferuje też specjalną płytkę zawierającą procesory W65C02 i W65C816 (ten drugi stanowi 16-bitową wersję modelu 6502). Twórca OtterX stworzył także emulator układu YM2151 w formie małej płytki z układem FPGA Lattice ICE40, którego można używać jako przetwornik DAC: zamiast oryginalnego i nieprodukowanego już YM3012 albo typowego układu DAC I²S. Sam ten moduł jest na tyle ciekawy, że mam ochotę go zakupić (40 dolarów z wysyłką do Polski) i wykorzystać w jakimś projekcie audio.

Wspomniałem wcześniej o DEC PDP-8. Czytelnik może zwrócić uwagę, że istnieje projekt PiPDP-11, czyli emulator późniejszego komputera firmy Digital Equipment Corporation, z własną płytką udającą panel kontrolny i pracujący na Raspberry Pi. Emulowanie komputera 16-bitowego o wydajności 0,2–1 MIPS i pojemności pamięci RAM – przeważnie – 64 kB...1 MB na 64-bitowym komputerze (mającym 512 MB pamięci RAM i wydajność powyżej 1000 MIPS) to nie jest żadne wyzwanie. Dlatego takie projekty całkowicie mnie nie interesują. Generalnie nie interesuje mnie żaden projekt z komponentem sprzętowym, w którym emulacja programowa jest realizowana na układzie trzy rzędy wielkości wydajniejszym od oryginału. Wyzwaniem byłoby zrealizowanie takiej emulacji na układzie 16- lub 32-bitowym, mającym poniżej 50 MIPS. Moim skromnym marzeniem jest stworzyć emulator PDP-8 na układzie PIC24/dsPIC oraz sprzętową emulację komputera Odra 1305/SKOK na układach CPLD lub małych FPGA. Definitywnie do tego drugiego rozwiązania przydałby się stosowny kurs.

Skoro było o starych komputerach, to warto wspomnieć też o technologiach powiązanych, które także już odeszły w niepamięć. Kto z Czytelników pamięta czasy, gdy programy i gry były na dostępne dyskietkach? Kto pamięta kartridże do różnych konsol, od Atari 2600 po GameBoya? Albo używanie kaset magnetofonowych do przechowywania gier? Ponoć w latach 80. jeden z pracowników Polskiego Radia późną nocą puszczał taśmy z programami, by słuchacze mogli sobie zgrać kopie z radia, ale osobiście myślę, że może to być kolejna miejska legenda, jak ta o czarnej Wołdze. Tak czy inaczej nośniki taśmowe z programami czy gramami zniknęły wraz z erą ośmiobitowych komputerów i nawet retromaniacy używają emulatorów na bazie kart SD lub CF. ZUS ostatni raz zamawiał dyskietki 3,5” 1,44 MB w 2008 roku, za to w imponującej ilości 130 tysięcy sztuk, dzięki czemu instytucja stała się na krótko pośmięskiem w polskim Internecie. Dobrze, że nie potrzebowali dyskietek 5,25” czy 8”. Spośród magnetycznych nośników danych tylko twarde dyski mają się naprawdę dobrze, a profesjonalści używają taśm LTO o pojemnościach dochodzących do 30 TB (dla LTO10). Przydałby mi się skromniejszy streamer LTO3 albo LTO4, ale to urządzenia typowo serwerowe, które trudno skłonić do pracy ze zwykłym komputerem domowym.

Kilka lat temu przeglądałem znany portal aukcyjny, a konkretnie dział używanych oscyloskopów analogowych. Uderzyły mnie wtedy dwie rzeczy: mocno zawyżone ceny sprzętów starszych ode mnie o dobre 10...20 lat oraz jedna, szczególna oferta oscyloskopu analogowego, opisanego jako „vintage, retro, loft”. Stan techniczny? Nieznany, ale się włącza. Parametry? Sprzedawca nie wie. Ale jak ładnie będzie wyglądać na półce. Cena? Porównywalna z dwukanałowym oscyloskopem cyfrowym z dolnej, ale użytecznej półki. Obecnie takie oferty trafiają się dużo rzadziej, ale wciąż nie brakuje oscyloskopów analogowych o przeciętnych parametrach i cenach porównywalnych z nowym DSO o 2...3 razy szerszym paśmie, takich marek jak Hantek, Owon czy UNI-T. Swoją drogą, wśród innych produktów vintage/retro/loft znalazłem kiedyś kompletny automat telefoniczny na żetony, używany przez TP S.A. jeszcze we wczesnych

latach dwutysięcznych, w czasach mojego liceum. Co ciekawe, taki automat pozwalał na dłuższą rozmowę na jednym żetonie (3 minuty to był standard), trzeba było jedynie w odpowiednim momencie uderzyć kolanem w spód automatu, dzięki czemu żeton przelatował przez czujnik w mechanizmie wrzutowym dwa razy: raz do góry i raz na dół. Powiedziałem o tym koledze z klasy, który metodę przetestował wieczorową porą w internacie. Metoda zadziałała, ale wraz z żetonem w mechanizmie podskoczyły wszystkie żetony w skarbonce. Huk był znaczny.

Teraz budek telefonicznych już nie ma, zniknęły nawet niebieskie i srebrne Urmetry oraz tzw. „Jajka”, które można było otworzyć wyciągając uszczelkę pomiędzy połówek obudowy i trochę nią manipulując. Nie, żeby to komuś coś dało bez sporej ilości wiedzy i odpowiednich narzędzi, jak chociażby klonu karty serwisowej. W roku 2001 czytałem o tym, jak grupa phreakerska (phreaking – hacking systemów telefonicznych, głównie budek, celem dzwonięcia za darmo, zwłaszcza międzymiastowo; dziś praktycznie zapomniane zjawisko, nie licząc kilku krajów rozwijających się) „Urmet Developers” stworzyła nagrywarkę kart telefonicznych na bazie mechanizmu czytnika skradzionego z budki, a także programowy emulator kart chipowych na mikrokontrolerze Atmela. Niestety dla nas, zwykłych nastolatków zainteresowanych elektroniką i darmowymi rozmowami, te narzędzia i kody źródłowe nie były dostępne. Młodzi Czytelnicy nie pamiętają raczej tych czasów, gdy rozmowa telefoniczna na ulicy wymagała stania przy ścianie lub słupie z budką na żetony lub karty, a posiadanie jakiegokolwiek telefonu komórkowego było zarezerwowane dla ludzi biznesu i to tych dobrze zarabiających. Teraz takie budki można znaleźć tylko w muzeach, albo czasem na aukcji, by móc urządzić sobie loft w stylu vintage. Do zestawu brakuje jeszcze saturatora, telewizora „Rubin” albo „Jowisz” i pralki „Frani” z maglem. Ale kto wie, może wróci moda na te klasyczne urządzenia domowe. Dwie dekady temu była widoczna moda na pralki typu „Frania” od różnych producentów, gdyż świetnie pasowały do skromnych lokali i budżetów typowych studentów. Teraz w modzie są pralki i pralko-suszarki z „AI”, czyli czymś, co kiedyś nazywaliśmy programatorem. O, kolejna zapomniana technologia: programator, w którym gałka połączona była z bębniem, na którym znajdowały się wycięcia, a kilka krzywek przesuwano się po tych wycięciach, przełączając pracę silnika, pompy wody brudnej czy zaworu wody ciepłej, a także innych elementów. Z tyłu zaś niewielki silnik synchroniczny obracał bęben (i zarazem gałkę) za pomocą serii przekładni, by jeden program na $\frac{1}{8}$ obrotu trwał godzinę lub dłużej. Takie programatory były jeszcze w użyciu we wczesnych latach 90., ale dość szybko wyparły je elektroniczne sterowniki na mikrokontrolerach, w tym 8051 – układzie, który w różnych formach i wariantach wciąż jest bardzo popularny i stosowany w elektronice (także konsumenckiej).

Model biznesowy na „retro”

Duże korporacje od jakiegoś czasu wypuszczają na rynek produkty konsumenckie, które – jeśli same nie są reinkarnacją zapomnianych urządzeń – to przynajmniej naśladują je wyglądem. I tak nic nie stoi na przeszkodzie, by kupić sobie radio DAB+, które stylem przypomina radia lampowe lat 50. Nowe magnetofony z marnymi mechanizmami, ale wyglądem a’la lata 80. i 90.? Proszę bardzo. Ba, można kupić nawet telefon komórkowy na wzór legendarnej, pancernej Nokii 3310. O retrokonsolach i retrokomputerach już wspominałem wcześniej. Tylko czekać, aż ktoś wypuści DSO z wyglądem starego oscyloskopu analogowego. Na rynku wciąż są dostępne multimetry analogowe czy lutownice transformatorowe dla tych elektroników, którzy zatrzymali się w latach 80. ubiegłego wieku. Polecanie początkującym lutownicy innego typu niż oporowa z kontrolą temperatury (z lub bez stacji hot-air) uważam obecnie za propagowanie szczególnie patologicznych zachowań, co moim zdaniem powinno być ścigane prawem...

Warto też pochylić się nad kreatywnymi „biznesmenami”, którzy na znanych portalach aukcyjnych i ogłoszeniowych próbują przehandlować wygrzebane ze strychów i piwnic, kupione na pchlich targach czy znalezione na śmietniku sprzęty starej Unity bądź wyroby samochodopodobne FSO czy FSM jako jako pamiątki z przeszłości i przykłady piękna polskiej inżynierii. Jeśli Czytelnik jest miłośnikiem makaronu, to zapraszam do zwiedzania wnętrza radiomagnetofonów Unity, w niektórych jest tego tyle, że nic, tylko otwierać włoską restaurację. Dziesiątki cienkich, kolorowych przewodów, łączących najróżniejsze punkty na płytkach drukowanych, wykonanych z (często) wątpliwej jakości laminatu. Jeśli mamy szczęście, znajdziemy też gniazda i wtyczki będące sowiecką odpowiedzią na goldpiny. Jakość ich była tak niska, że do urządzeń trafiały już fabrycznie zaśniedziałe, więc nawet ówczesni serwisanci ich nienawidzili. Ale sprytny Polak bez problemu wciśnie naiwnemu klientowi produkt niesprawny twierdząc, że sam nie sprawdzał i sprzedaje sprzęt taki, jaki jest. Caveat emptor, innymi słowy.

To na co może się skusić współczesny miłośnik retro? O absurdalnych cenach gramofonów już wspominałem. Za trochę ponad 3500 złotych (w chwili pisania tego artykułu) można nabyć amplifier Unitra Elizabeth Stereo DST-202. Ten reprezentant piękna polskiej inżynierii z 1973 roku ma wątpliwą użyteczność: na falach długich i tak jest słyszalna „sieczka” od przetwornic, a UKF może wymagać przestrojenia na zakres 88...108 MHz. Wzmacniacz mocy 2×10 W też nie wprawi w zachwyt, a może wymagać naprawy po tylu latach. Czytelnik lepiej zrobi, jeżeli zbuduje własny wzmacniacz i przedwzmacniacz. Projektowanie takiego układu od zera to niezła przygoda w świecie elektroniki analogowej, ale dla bardziej leniwych Czytelników w ofercie AVT są dziesiątki, jeśli nie setki kitów audio. Sam mam do skończenia taki zestaw oparty o 2×LM3886T, czyli popularny wśród budżetowych audiofilów wariant wzmacniacza GainClone. Wspominałem już o problematycznych magnetofonach z epoki. Na tym też się zarabia. Pamiętam te czasy, gdy stare „Jowity” leżały stertami pod śmietnikami, obok stosów „Rubinów” i „Jowisz” wypartych z domów przez tanie telewizory i radiomagnetofony z zachodu. Oczywiście nie tylko uznanych marek, ale też ich podróbek, kupionych tanio na bazarach: telewizorów Panasonic, walkmanów Somy, czy chińskich klonów Pegazusa (który sam był klonem NESa z Tajwanu), sprzedawanych jako PolyStation. Swoją drogą nawet te tanie kopie są popularne na rynku retro, co dowodzi, iż na każdym szmelcu można zarobić. Vide sprzedawane niegdyś 200 złotych „duże fiaty” z FSO potrafią kosztować 10...20 tysięcy złotych, a rekordowa cena w ogłoszeniu, które widziałem, wynosiła 34 tysiące złotych. Kto by chciał jeździć tym motoryzacyjnym troglodytą, nie mam pojęcia.

Jaką przeszłość przyniesie przyszłość?

Retromania dosięgła lat 80. kilka lat temu, wraz ze wzrostem zainteresowania komputerami ośmiobitowymi z epoki. Spodziewam się, że następne będą komputery z wczesnych i późnych lat 90., czyli modele od 486 do Pentium II i pierwszych akceleratorów 3D. Remake’i i remastery gier z lat dwutysięcznych pojawiają się już od paru lat. W branży audio czeka nas powrót płyt CD, discmanów, a później odtwarzaczy MP3. Ciekawe, czy powrócą też odtwarzacze MD (MiniDisc). Wśród fanów ruchomych obrazów renesans przeżyje taśma VHS, już teraz będąca cichą bohaterką twórców tzw. „analog horror”, gatunku filmowego popularnego na YouTube. Ciężko jest komputerowo uzyskać realistyczny wygląd zdegradowanego nagrania, wykonanego kamkorderem na taśmie VHS-C czy Hi8. Nie będę jednak namawiać Czytelników do zbierania starego sprzętu, bo mimo wszystko mam nadzieję, że trend retromanii umrze czym prędzej, a nam pozostanie tylko podziwiać nowe urządzenia wyglądające jak starocie z czasów dawno minionych.

Paweł Kowalczyk, EP

Programowanie w środowisku MicroPython (5)

Wyświetlacz OLED



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

Monochromatyczne wyświetlacze OLED o niskiej rozdzielczości są ciekawą i tanią alternatywą dla 7-segmentowych wyświetlaczy LED i tekstowych paneli LCD. Pozwalają przy tym stworzyć dużo ciekawszy i intuicyjny interfejs użytkownika, zachowując prostotę i niski koszt.

W tym odcinku kursu nauczymy się obsługiwać wyświetlacz OLED o przekątnej 2,42 cala i rozdzielczości 128×64 pikseli z wbudowanym sterownikiem SSD1309 firmy Solomon Systech. Takie wyświetlacze produkuje chińska firma Wisevision Optronics. Dostępne są w kolorach: białym, żółtym, zielonym i niebieskim. Istnieją także tańsze wyświetlacze o mniejszych rozmiarach i tej samej rozdzielczości. Na rynku można znaleźć podobne rozwiązania z nieco innymi sterownikami, takimi jak SSD1306 lub SH1106. Wszystkie te sterowniki są bardzo podobne i kod, który napiszemy dla SSD1309, powinien działać bez żadnych modyfikacji także z innymi sterownikami.

SSD1309 jest układem scalonym montowanym w technologii chip-on-glass co oznacza, że jest zintegrowany z szybą wyświetlacza, zza której wyprowadzony jest 24-żyłowy kabel. Należy go włożyć do gniazda FFC/FPC o rastrze 0,5 mm. Do uruchomienia wyświetlacza potrzebujemy dwóch napięć zasilających (3,3 V i 12 V) oraz kilku drobnych elementów pasywnych.

Istnieje kilka niedrogich płytek, które wszystkie niezbędne połączenia mają już gotowe do wykorzystania. Są to chińskie produkty typu „no name” – bez marki wytwórcy ani innych oznaczeń – dostępne w wielu sklepach. Mają one kilka charakterystycznych cech, dzięki którym można je rozróżnić. Są to: kolor soldermaski, liczba pinów konektora i jego położenie oraz typ interfejsu komunikacyjnego.

Na początek zdecydowanie odradzam zakup płytek z czarną soldermaską i 4-pinowym konektorem z interfejsem I²C. Płytki te mają jakiś problem z przetwornicą podwyższającą napięcie – w rezultacie konwerter bardzo nieprzyjemnie piszczy, co jest mocno irytujące, zwłaszcza w nocy.

Polecam natomiast płytki z niebieską soldermaską i 7-pinowym konektorem umieszczonym na górze lub po lewej stronie PCB. Przykład takiego modułu pokazuje **fotografia 1**. Mogą one pracować z interfejsem SPI lub I²C. Fabrycznie płytki są skonfigurowane tak, by pracowały z interfejsem SPI, ale przeróbka na I²C jest bardzo łatwa. Wystarczy do tego lutownica, cyna i kawałek drutu.

Na dolnej stronie płytki znajdują się wszystkie drobne elementy elektroniczne. Aby przerobić płytkę z interfejsu SPI na I²C, musimy najpierw wyjąć kabel wyświetlacza. W tym celu delikatnie przesuwamy szarą klamrę gniazda w kierunku środka płytki, po czym przewód możemy delikatnie wysunąć ze złącza. W pierwszej kolejności musimy odlutować rezystor R8 (0 Ω) i włutować go w miejsce R9. Pola lutownicze są na tyle duże, że zamiast rezystora możemy umieścić tam kropelkę cyny. Podobnie robimy w miejscu R11.

Kolejna możliwa modyfikacja jest niepotrzebna, jeżeli chcemy zasilać wyświetlacz ze źródła o napięciu 5 V. Na płytce jest umieszczony stabilizator LDO konwertujący 5 V na 3,3 V. Jeżeli chcemy zasilać płytkę bezpośrednio ze źródła 3,3 V, warto obejść stabilizator, lutując mały drucik zwierający wejście i wyjście regulatora (najlepiej



Zobacz więcej:

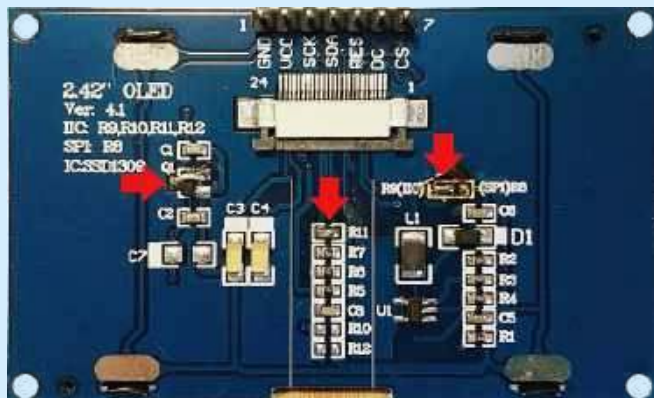
- Repozytorium kursu na GitHubie: <https://github.com/leonow32/micropython>
- Dokumentacja sterownika SSD1309: https://support.newhavendisplay.com/hc/en-us/article_attachments/4414433936535
- Dokumentacja wyświetlacza Wisevision Optronics OLED 128×64 2.42" yellow: https://www.lcsc.com/datasheet/lcsc_datasheet_2410121457_Newvisio-X242-2864KSUG01-C24_C5123572.pdf
- Dokumentacja klasy FrameBuffer: <https://docs.micropython.org/en/latest/library/framebuf.html>
- Algorytm Bresenham'a do rysowania linii prostych: https://en.wikipedia.org/wiki/Bresenham%27s_line_algorithm
- Dekoratory Viper i Native: https://docs.micropython.org/en/latest/reference/speed_python.html#the-native-code-emitter

uprzednio wylutowując stabilizator – przyp. red.). Wszystkie modyfikacje zaznaczono czerwonymi strzałkami na **fotografii 2**.

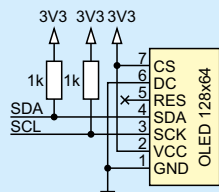
Możemy podłączyć płytkę z wyświetlaczem do ESP32-S3 w sposób pokazany na **rysunku 1**. Domyślny pin SDA to GPIO8,



Fotografia 1. Płytkę z wyświetlaczem OLED 128×64 o przekątnej 2,42 cala



Fotografia 2. Dolna strona płytki po wykonanych modyfikacjach



Rysunek 1. Sposób połączenia płytki ESP32 z wyświetlaczem

a w przypadku SCL – GPIO9, ale w konstruktorze klasy I²C można wskazać dowolne inne piny, jeżeli z jakiegoś powodu linie 8 i 9 chcemy wykorzystać w innym celu. Sterownik SSD1309 może używać dwóch adresów na magistrali I²C. Jeżeli pin DC połączymy z masą, wówczas sterownik będzie występował pod adresem 0x3C, a jeżeli do zasilania – adres zmieni się na 0x3D. Pin RES służy do resetowania kontrolera SSD1309. Na płytce znajduje się układ RC, który zainicjalizuje sterownik zaraz po włączeniu zasilania, więc wspomnianą linię możemy zostawić niepodłączoną.

Po podłączeniu wyświetlacza do ESP32 warto (w celach diagnostycznych) przeskanować wszystkie adresy na magistrali I²C przy pomocy skryptu `i2c_scan.py`, który opracowaliśmy w 3. odcinku kursu. Jeżeli na konsoli pojawi się komunikat:

Znalezione urządzenia: 3C

...to możemy przejść do kolejnego rozdziału.

Sterownik SSD1309

W niniejszym artykule omówimy wyłącznie komunikację z kontrolerem SSD1309 za pośrednictwem interfejsu I²C. Wymiana danych jest jednokierunkowa, za wyjątkiem bitów potwierdzeń ACK, którymi SSD1309 sygnalizuje odebranie danych z szyny I²C. Nie ma możliwości, by odczytać obraz z wyświetlacza ani jakiegokolwiek inne dane.

Pierwszym bajtem wysyłanym przez I²C jest adres urządzenia slave, czyli 0x3C lub 0x3D – w zależności od sposobu podłączenia pinu DC. Kolejny bajt to tzw. *control byte*. Może on przyjąć jedną z dwóch wartości:

- 0x80 – następny bajt będzie poleceniem do wykonania. Jeżeli chcemy przesłać więcej komend w jednej transmisji, to każdy bajt polecenia musimy poprzedzić bajtem 0x80.
- 0x40 – wszystkie kolejne bajty, aż do zakończenia transmisji, są danymi obrazu do wyświetlenia.

Sterownik SSD1309 obsługuje całkiem sporo różnorodnych poleceń. Omówimy tylko kilka z nich, a Czytelników zainteresowanych pozostałymi instrukcjami odsyłam do dokumentacji dostępnej pod adresem [2].

Wyświetlacz ma rozdzielczość 128 pikseli w poziomie i 64 w pionie. Piksele są monochromatyczne, czyli albo się świecą, albo pozostają wygaszone.

Ekran podzielono na osiem pasów o wysokości 8 pikseli. W dokumentacji SSD1309 pasy te nazywane są stronami (page). Wysokość 8 pikseli nie jest przypadkowa – dokładnie tyle bitów zawiera jeden bajt, czyli najmniejsza jednostka danych, jaką możemy przesłać przez interfejs I²C. A zatem przysyłając jeden bajt ustawiamy stan kolumny, złożonej z 8 pikseli. Cursor inkrementuje się automatycznie i przesuwają w prawo o jedną pozycję po odebraniu każdego bajtu danych.

Kursor możemy ustawić na dowolną stronę i dowolną współrzędną x. Możemy także zmieniać zakres automatycznej inkrementacji kursora. W praktycznych rozwiązaniach kursor umieszczamy na współrzędnej (0,0), a zakres jego przemieszczania obejmuje cały wyświetlacz. W taki sposób możemy odświeżyć cały obszar wyświetlacza w pojedynczej transakcji I²C. Skoro mamy 128 kolumn w 8 stronach, to przesłać musimy 1024 bajty.

Sterownik SSD1309 musimy zainicjalizować po każdym włączeniu zasilania lub zresetowaniu, przysyłając do niego szereg różnych poleceń. Często producenci wyświetlaczy podają w dokumentacji sposób, w jaki sterownik musi być skonfigurowany, aby działał poprawnie. Ustawiać można różne tryby multipleksacji, działanie przetwornicy podwyższającej napięcie, częstotliwość odświeżania i wiele innych parametrów. Te rzeczy to „wiedza tajemna”. Konieczność ustawiania takich parametrów wynika z faktu, iż wielu producentów stosuje sterownik SSD1309 z wyświetlaczami o różnych rozmiarach i parametrach, a kontroler nie ma wbudowanej pamięci Flash, przez co nie może być skonfigurowany przez producenta wyświetlacza. Przykład sekwencji inicjalizacyjnej znajdujemy w dokumentacji firmy Wisedvision Optronics pod adresem [3].

Dla nas najbardziej istotne będą instrukcje odpowiedzialne za regulację kontrastu i możliwość obrócenia obrazu o 180°. Omówimy to dokładniej analizując kod programu.

Biblioteka FrameBuffer

Do generowania grafiki wykorzystamy moduł **FrameBuffer**, który jest standardowo wbudowany w implementację MicroPythona na ESP32. Dokładny opis wszystkich możliwości tego modułu znajduje się pod adresem [4].

Biblioteka dostarcza różnych funkcji pozwalających narysować pojedyncze piksele, linie, prostokąty, okręgi i inne figury geometryczne. Możemy tworzyć napisy (choć w standardzie dostępny jest tylko jeden font) albo skorzystać z przygotowanych wcześniej plików graficznych.

Na początku programu musimy utworzyć bufor, w którym biblioteka będzie rysować wszystkie kształty. Po zakończeniu rysowania klatki obrazu przesyłamy całą zawartość bufora do wyświetlacza. W tym celu musimy sami opracować jakąś funkcję, która odczyta gotowy bufor w pamięci RAM, prześle go do wyświetlacza przez interfejs komunikacyjny i, jeżeli to potrzebne, wyśle również określone polecenia.

Omówimy teraz najbardziej przydatne funkcje z klasy **FrameBuffer**. Zaczniemy od jej konstruktora, który przyjmuje cztery argumenty – są to:

- obiekt typu `bytearray`, w którym znajdzie się bufor roboczy obrazu. Ważne, by jego rozmiar był dopasowany do rozdzielczości wyświetlacza i sposobu kodowania pikseli. W naszym wyświetlaczu OLED na jeden bajt przypada 8 pikseli. W przypadku wyświetlacza kolorowego w formacie RGB565 każdy piksel korzysta z 2 bajtów,
- szerokość w pikselach,
- wysokość w pikselach,
- format obrazu. W przypadku wyświetlaczy monochromatycznych istnieje kilka możliwości, ale dla SSD1309 właściwą opcją jest `MONO_VLSB`. W przypadku wyświetlaczy kolorowych najczęściej stosuje się format `RGB565`.

Instancję klasy **FrameBuffer** możemy utworzyć w następujący sposób:

```
WIDTH = const(128)
HEIGHT = const(64)
array = bytearray(WIDTH * HEIGHT // 8)
buffer = FrameBuffer(array, WIDTH, HEIGHT, MONO_VLSB)
```

Uzyskujemy w ten sposób dostęp do różnych metod z klasy **FrameBuffer**. Pierwszą z nich jest **fill**, która wypełnia cały wyświetlacz jednolitym kolorem. W przypadku wyświetlaczy monochromatycznych mamy tylko dwa kolory – 0 oznacza piksel wygaszony, a 1 to piksel zaświecony.

```
buffer.fill(0) # Wszystkie piksele czarne
buffer.fill(1) # Wszystkie piksele świecą
```

Oczywiście możemy kontrolować także każdy piksel osobno. W tym celu korzystamy z metody **pixel**, do której przekazujemy współrzędne x, y oraz żądany kolor.

```
buffer.pixel(x, y, color)
```

Do rysowania linii mamy aż trzy metody. Pierwsza z nich to **line**, rysująca linię prostą od punktu (x1, y1) do punktu (x2, y2) przy użyciu algorytmu Bresenhama, który został dokładnie omówiony pod adresem [5]. Jest on uniwersalny, ale w przypadku linii pionowych okazuje się nieoptymalny – dużo szybciej linię narysujemy stosując zwyczajną pętlę, która koloruje piksele po kolei. W tym celu autorzy MicroPythona dodali metody **hline** (do rysowania linii poziomych) oraz **vline** (do pionowych). W obu wspomnianych metodach podajemy współrzędne początku linii, jej długość oraz kolor.

```
buffer.line(x1, y1, x2, y2, color)
buffer.hline(x, y, w, color)
buffer.vline(x, y, h, color)
```

W bibliotece mamy dostępny tylko jeden font o wysokości 8 pikseli i stałej szerokości znaków, równej także 8 pikseli. Font jest mało estetyczny, a litery – niesymetryczne. Nie ma możliwości wyrównania tekstu do prawej ani do środka. Dlatego w dalszej części artykułu poznamy sposób pozwalający na rozwiązanie tego problemu i zastosowanie własnych fontów.

Aby wyświetlić napis wbudowanym fontem we **FrameBuffer**, należy skorzystać z metody **text**.

```
buffer.text("Napis do wyświetlenia", x, y, color)
```

Wyświetlenie bitmapy jest dość trudne. Trzeba ją najpierw przekonwertować na obiekt typu **bytearray**, w którym wszystkie bity są ustawione w takim samym formacie jak **bytearray** bufora obrazu. Następnie tworzymy drugi obiekt klasy **FrameBuffer**, a potem z bufora obrazu wywołujemy metodę **blit**. Pierwszym argumentem jest utworzony wcześniej **FrameBuffer** bitmapy, a następne parametry to współrzędne x i y, w których ma znaleźć się lewy górny narożnik bitmapy. W jaki sposób przekształcić plik graficzny na **bytearray** – o tym będzie mowa w dalszej części artykułu.

```
moja_bitmapa = framebuffer.FrameBuffer(
    bytearray(b'...dane bitmapy...'),
    szerokość, wysokość, framebuffer.MONO_VLSB)

buffer.blit(moja_bitmapa, x, y)
```

Klasa SSD1309

Moduł **FrameBuffer** nie ma żadnego powiązania ze sprzętem – generuje on tylko grafikę w pamięci RAM, którą sami musimy przesłać do wyświetlacza. Możemy do tego tematu podejść na dwa sposoby.

Pierwszym jest utworzenie instancji klasy **FrameBuffer**, tak jak to zrobiliśmy w poprzednim rozdziale. „Obok” niej tworzymy osobną klasę odpowiadającą za komunikację z wyświetlaczem. Chcąc wyświetlić klatkę obrazu na wyświetlaczu, przekazujemy **bytearray** z **FrameBuffera** do funkcji transmitującej ów bufor przez I²C do wyświetlacza. Takie rozwiązanie jest oczywiście poprawne, ale mało wygodne, bo mamy dwie osobne klasy. Lepiej byłoby wszystko, co związane z wyświetlaczem, wrzucić do jednego worka.

Z pomocą przychodzi dziedziczenie. Możemy zrobić klasę **SSD1309** w taki sposób, że wszystkie metody klasy **FrameBuffer** będą działać tak, jakby one były metodami klasy **SSD1309**. Jedyne, co w klasie **SSD1309** musimy zrobić, to obsługa sprzętu. Oprócz tego dodamy kilka metod umożliwiających bardziej zaawansowane wyświetlanie napisów różnymi fontami. Nasza biblioteka będzie miała możliwość obrócenia obrazu o 180°, a także zmiany adresu sterownika na magistrali I²C. Ciekawą funkcjonalnością będzie metoda wyświetlająca bufor obrazu jako ASCII-art w konsoli, co pozwoli przetestować kod wszystkim tym Czytelnikom, którzy nie posiadają akurat pod ręką wyświetlacza ze sterownikiem **SSD1309**. Weźmy pod lupę kod modułu **SSD1309** pokazany na **listingu 1**.

Zaczynamy od zaimportowania modułu **framebuf** w linii 1, po czym tworzymy dwie stałe, odpowiadające za szerokość i wysokość wyświetlacza. W naszym przypadku jest to 128 pikseli (szerokość) i 64 (wysokość). Aby nasza biblioteka była bardziej uniwersalna i obsługiwała mniejsze wyświetlacze, można pomyśleć o ustawianiu tych wartości poprzez konstruktor klasy. Przypomnijmy, że stałe **const()** to twórcy dostępny tylko w MicroPythonie, który pozwala zaoszczędzić pamięć RAM. Nie znajdziemy nic podobnego w klasycznym Pythonie.

W linii 3 rozpoczynamy klasę o nazwie **SSD1309**. W nawiasach po jej nazwie podajemy nazwę klasy, z której mają zostać odziedziczone wszystkie metody – w naszym przypadku jest to **FrameBuffer** z modułu **framebuf**.

Klasę rozpoczynamy metodą specjalną **__init__**, czyli konstruktorem. Przekazujemy do niego kilka argumentów, które zostaną następnie zapisane w postaci zmiennych wewnętrznych klasy, między innymi: **i2c** oraz **address**, czyli instancja interfejsu I²C oraz adresem sterownika **SSD1309**. Dzięki takiemu rozwiązaniu możemy wykorzystać tę samą klasę, aby sterować więcej niż jednym wyświetlaczem – obojętnie, czy ekrany będą podłączone do wielu różnych interfejsów, czy do jednego, ale na różnych adresach.

Argument **rotate** może przyjmować wartości **True** lub **False**, a jego celem jest odwrócenie (lub nie) obrazu na wyświetlaczu o 180°. Wykorzystujemy tutaj sprzętowe możliwości sterownika. Stan tego argumentu modyfikuje niektóre instrukcje w sekcji inicjalizacyjnej, która jest przesyłana do sterownika kilka linijek niżej.

W linii 5 tworzymy zmienną **array**. Przechowujemy w niej bufor obrazu, będący obiektem typu **bytearray**. W nawiasach musimy podać rozmiar bufora – musi on pomieścić wszystkie piksele. Mnożymy więc szerokość i wysokość, a następnie dzielimy wynik przez 8, ponieważ w jednym bajcie mieści się 8 pikseli. Małe przypomnienie – w języku Python operator / zwraca wynik dzielenia jako liczbę zmiennoprzecinkową, nawet jeżeli po przecinku jest zero. Operator // zwraca natomiast zawsze liczbę całkowitą – a taki musi być przecież rozmiar bufora.

Klasa **FrameBuffer** też ma swój konstruktor, który musimy wywołać, aby ją odpowiednio skonfigurować. By uzyskać dostęp

```

# Plik ssd1309.py

import framebuffer # 1

WIDTH = const(128) # 2
HEIGHT = const(64)

class SSD1309(framebuf.FrameBuffer): # 3

    @micropython.native
    def __init__(self, i2c, rotate=False, address=0x3C): # 4
        self.i2c = i2c
        self.address = address
        self.array = bytearray(WIDTH * HEIGHT // 8) # 5
        super().__init__(self.array, WIDTH, HEIGHT, framebuffer.MONO_VLSB) # 6

        config = ( # 7
            0xAE, # Display disable
            0x20, 0x00, # Set memory addressing mode
                    # to horizontal addressing mode
            0x40, # Set display start line to 0
            0xA0 if rotate else 0xA1, # Set segment remap
            0xA8, 0x3F, # Set multiplex ratio to 63
            0xC0 if rotate else 0xC8, # Set COM scan direction
            0xD3, 0x00, # Set display offset to 0
            0xDA, 0x12, # Set COM pins hardware config
                    # to enable COM left/right remap
            0xD5, 0x80, # Set clock and oscillator frequency
            0xD9, 0xF1, # Set pre-charge period
            0xDB, 0x3C, # Set VCOMH to max
            0x81, 0xFF, # Set contrast to 255 (max)
            0xA4, # Use image in GDDRAM memory
            0xA6, # Display not inverted
            0xAF, # Display enable
        )

        for cmd in config: # 8
            self.write_cmd(cmd)

    @micropython.viper
    def write_cmd(self, cmd: int): # 9
        self.i2c.writeto(self.address, bytes([0x80, cmd]))

    @micropython.viper
    def display_on(self): # 10
        self.write_cmd(0xAE)

    @micropython.viper
    def display_off(self): # 11
        self.write_cmd(0xAE)

    @micropython.viper
    def contrast(self, value): # 12
        self.write_cmd(0x81)
        self.write_cmd(value)

    @micropython.viper
    def refresh(self): # 13
        for cmd in (0x21, 0x00, 0x7F, 0x22, 0x00, 0x07): # 14
            self.write_cmd(cmd)

        self.i2c.writevto(self.address, (b"\x40", self.array)) # 15

    @micropython.viper
    def simulate(self): # 16
        for y in range(HEIGHT):
            print(f"{y}\t", end="")
            for x in range(WIDTH):
                bit = 1 << (y % 8)
                byte = int(self.array[(y // 8) * WIDTH + x])
                pixel = "#" if byte & bit else "."
                print(pixel, end="")
            print("")

    @micropython.native
    def print_char(self, font, char, x, y, color=1): # 17
        try:
            bitmap = font[ord(char)]
        except:
            bitmap = font[0]
            print(f"Char {char} doesn't exist in font")

        width = bitmap[0]
        height = bitmap[1]
        space = bitmap[2]

        if color:
            buffer = framebuffer.FrameBuffer(bitmap[3:], width, height, 0)
        else:
            negative_bitmap = bitmap[:3]
            for i in range(3, len(bitmap)):
                negative_bitmap[i] = ~negative_bitmap[i]
            buffer = framebuffer.FrameBuffer(negative_bitmap[3:], width, height, 0)
            self.rect(x-space, y, space, height, 1, True)
            self.rect(x+width, y, space, height, 1, True)

        self.blit(buffer, x, y)
        return width + space

    @micropython.native
    def print_text(self, font, text, x, y, align="L", color=1): # 18
        width = self.get_text_width(font, text)

        if align == "R":
            x = WIDTH - width
        elif align == "C":
            x = WIDTH // 2 - width // 2
        elif align == "l":
            pass

```

Listing 1. Kod pliku ssd1309.py

do niego musimy posłużyć się funkcją specjalną **super** (linia 6). W ten sposób możemy wywołać konstruktor **__init__** klasy nadrzędnej i skonfigurować ją, a dokładniej – przekazać do niej bufor obrazu utworzony linię wcześniej, a także szerokość, wysokość i format obrazu.

Następnym krokiem jest stworzenie sekwencji inicjalizacyjnej. W linii 7 tworzymy zmienną **config** i zapisujemy do niej krotkę, w której znajdują się wszystkie polecenia do przesłania do wyświetlacza w celu jego inicjalizacji. Wykonujemy to w pętli **for** (linia 8), gdzie każdy element krotki przekazujemy do metody **write_cmd**, zdefiniowanej w linii 9. Linia ta jest bardzo prosta i ogranicza się do wywołania metody **writeto**, należącej do klasy **i2c**, a przekazanej przez konstruktor. Do metody przekazujemy adres sterownika na magistrali oraz polecenie, jakie ma zostać wysłane. Zgodnie z notą katalogową sterownika, polecenie musimy poprzedzić bajtem kontrolnym o wartości 0x80. W tym celu ów bajt kontrolny – oraz bajt polecenia – obejmujemy nawiasami kwadratowymi, a następnie taki twór zapisujemy jako **bytes**. Moglibyśmy także wykorzystać **bytearray**. Przypomnijmy, że różnica między nimi jest następująca: **bytes** to format tylko do odczytu, a **bytearray** – do odczytu i zapisu poszczególnych składowych. Nasza metoda nic nie modyfikuje, zatem użycie **bytes** będzie optymalne.

Trzy kolejne metody demonstrują praktyczne zastosowanie **write_cmd**. W liniach 10 i 11 mamy metody, które włączają i wyłączają wyświetlacz. Całość sprowadza się do wysłania 1-bajtowego polecenia. W linii 12 mamy metodę ustawiającą kontrast wyświetlacza. Przyjmuje ona wartość w zakresie od 0 do 255. Aby ustawić kontrast, najpierw wysyłamy bajt 0x81, a zaraz po nim kolejne polecenie ustawiające wartość kontrastu. Każdy z tych dwóch bajtów poprzedzony jest bajtem kontrolnym 0x80, który automatycznie dodaje metoda **write_cmd**.

Dochodzimy do metody **refresh** w linii 13. Jej zadaniem jest przesłanie bufora array, zawierającego obraz, do sterownika SSD1309. Transmisja składa się z dwóch części. Najpierw musimy ustalić obszar obrazu, jaki chcemy zaktualizować. Wykonujemy to w sposób podobny do sekwencji inicjalizacyjnej, ale ponieważ musimy przesłać mniej poleceń, trochę to uprościmy. W linii 14 rozpoczynamy pętlę **for** iterującą po krotce, którą tworzymy w tej samej linii, co pętlę. Wewnątrz krotki znajduje się sześć bajtów – ich celem jest wskazanie prostokątnego obszaru na wyświetlaczu, który ma zostać zaktualizowany. Znaczenie kolejnych bajtów jest następujące:

1. 0x21 – dwa kolejne bajty to zakres, w którym może się zmieniać współrzędna x,
2. 0x00 – początek zakresu od współrzędnej x0=0,
3. 0x7F – koniec zakresu na współrzędnej x1=127,
4. 0x22 – dwa kolejne bajty to zakres, w którym może się zmieniać numer strony,
5. 0x00 – początek zakresu od strony 0,
6. 0x07 – koniec zakresu na stronie 7.

W ten sposób informujemy kontroler SSD1309, że chcemy zaktualizować cały obszar wyświetlacza o rozdzielczości 128×64, składający się z 8 stron o wysokości 8 pikseli każda.

Następnie, w linii 15, korzystamy z metody **wri-
tevt** z klasy **i2c**, która w jednej transakcji przesy-
ła podane „wektory”. Poprzez słowo *wektor* autorzy
MicroPythona rozumieją zmienne zawierające ciągi
bajtów. W naszym programie pierwszym takim ciągiem
jest pojedynczy bajt 0x40, który informuje sterownik,
że wszystkie kolejne bajty są grafiką do wyświetle-
nia. Drugim ciągiem jest oczywiście bufor **array**.

Metoda **simulate** (linia 16) ma za zadanie wyświet-
lić bufor obrazu na konsoli. Piksel widoczny oznacza-
ny jest znakiem #, a niewidoczny będzie reprezentowa-
ny jako kropka. W ten sposób na konsoli powstaje coś
à la ASCII art

Pozostało jeszcze do omówienia kilka metod, ale
na razie je pominiemy. Służą one do generowania napi-
sów z wykorzystaniem różnych fontów, co nie jest stan-
dardową funkcjonalnością MicroPythona. Wrócimy
do nich w dalszej części artykułu.

Na początku każdej metody znajdują się dziwne napi-
sy, takie jak **@micropython.viper** albo **@micropython.
native**. Są to tzw. dekoratory poprawiające wydajność kodu.
Funkcje oznaczone takimi dekoratorami muszą być napisane
w dość specyficzny sposób, ale za to wykonują się szybciej i zaj-
mują mniej pamięci.

@micropython.native powoduje, że kompilator Pythona tworzy
kod zawierający instrukcje wykonywane bezpośrednio przez pro-
cesor, a nie interpreter. W takim przypadku nie można używać blo-
ków **with**, generatorów, a ponadto są też pewne ograniczenia w ob-
słudze wyjątków.

@micropython.viper generuje kod jeszcze bardziej kompakto-
wy, ale ma to jeszcze wyższą cenę. Wszystkie argumenty funkcji
muszą mieć zdefiniowany typ, a wartość zwracana przez funk-
cję również musi mieć typ znany jeszcze zanim funkcja zacznie
się wykonywać. Możliwe typy to **int** (32-bitowa liczba całko-
wita ze znakiem), **uint** (liczba całkowita bez znaku), **ptr8**, **ptr16**,
ptr32 (wskaźniki na liczby odpowiednio, 8, 16 oraz 32-bitowe),
a także **ptr** (wskaźnik na dowolny obiekt). Poniżej można zobaczyć
przykład funkcji, która pobiera argument **arg** typu **int** i zwraca
wartość typu **uint**. Warto dodać, że funkcje oznaczone dekorato-
rem **@micropython.viper** nie mogą mieć argumentów z wartościami
domyślnymi.

@micropython.viper

```
def funkcja(self, arg: int) -> uint:
```

Więcej szczegółów na temat **native** i **viper** znajdziesz na stronie
pod adresem [6].

Testujemy!

Czas najwyższy wygene-
rować jakąś grafikę i poka-
zać ją na wyświetlaczu. Płytkę
z wyświetlaczem podłącz
tak, jak pokazuje to fotogra-
fia 2. Sygnał SCL należy do-
prowadzić do pinu GPIO 9
mikrokontrolera ESP32-S3,
a SDA – do GPIO 8. Ewentualnie
możesz użyć też innych pi-
nów. Zapisz plik **ssd1309.py**
w pamięci ESP32-S3,
a następnie otwórz nowy skrypt
i umieść w nim taki kod poka-
zany na listingu 2.

```
x = x - width + 1
elif align == "c":
    x = x - width//2

for char in text:
    x += self.print_char(font, char, x, y, color)

@micropython.native
def get_text_width(self, font, text):
    total = 0
    last_char_space = 0
    for char in text:
        bitmap = font.get(ord(char), font[0])
        total += bitmap[0]
        total += bitmap[2]
        last_char_space = bitmap[2]

    return total - last_char_space

if __name__ == "__main__":
    i2c = I2C(0)
    print(i2c)
    display = SSD1309(i2c)

    display.rect(0, 0, 128, 64, 1)
    display.text("abcdefghijklm", 1, 2, 1)
    display.text("nopqrstuvwxyz", 1, 10, 1)
    display.refresh()
    display.simulate()
```

Listing 1. Kod pliku **ssd1309.py** – **cd**.

```
# Plik demo_simple.py
from machine import Pin, I2C
import ssd1309

i2c = I2C(0)
print(i2c)

display = ssd1309.SSD1309(i2c)
# display = ssd1309.SSD1309(i2c, rotate=True)
# display = ssd1309.SSD1309(i2c, address=0x3D)

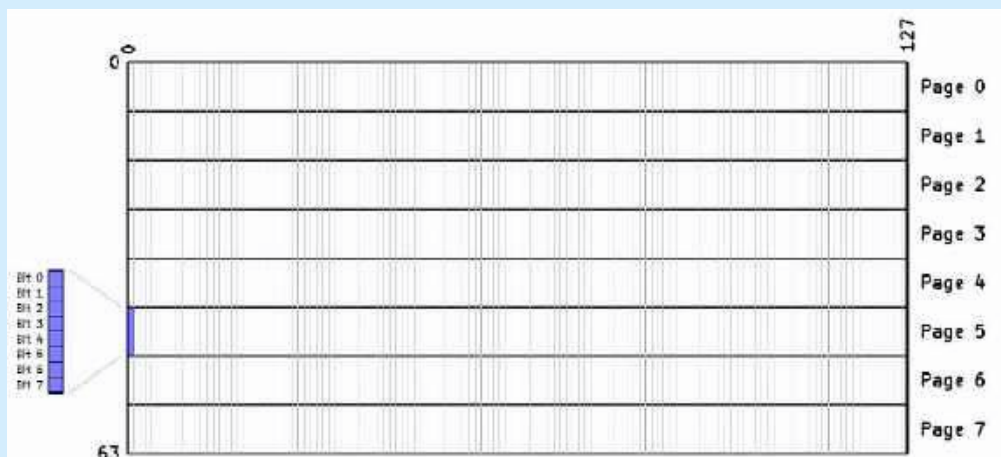
display.rect(0, 0, 128, 64, 1)
display.text('abcdefghijklm', 1, 2, 1)
display.text('nopqrstuvwxyz', 1, 10, 1)
display.refresh()
display.simulate()
```

Listing 2. Kod pliku **demo_simple.py**

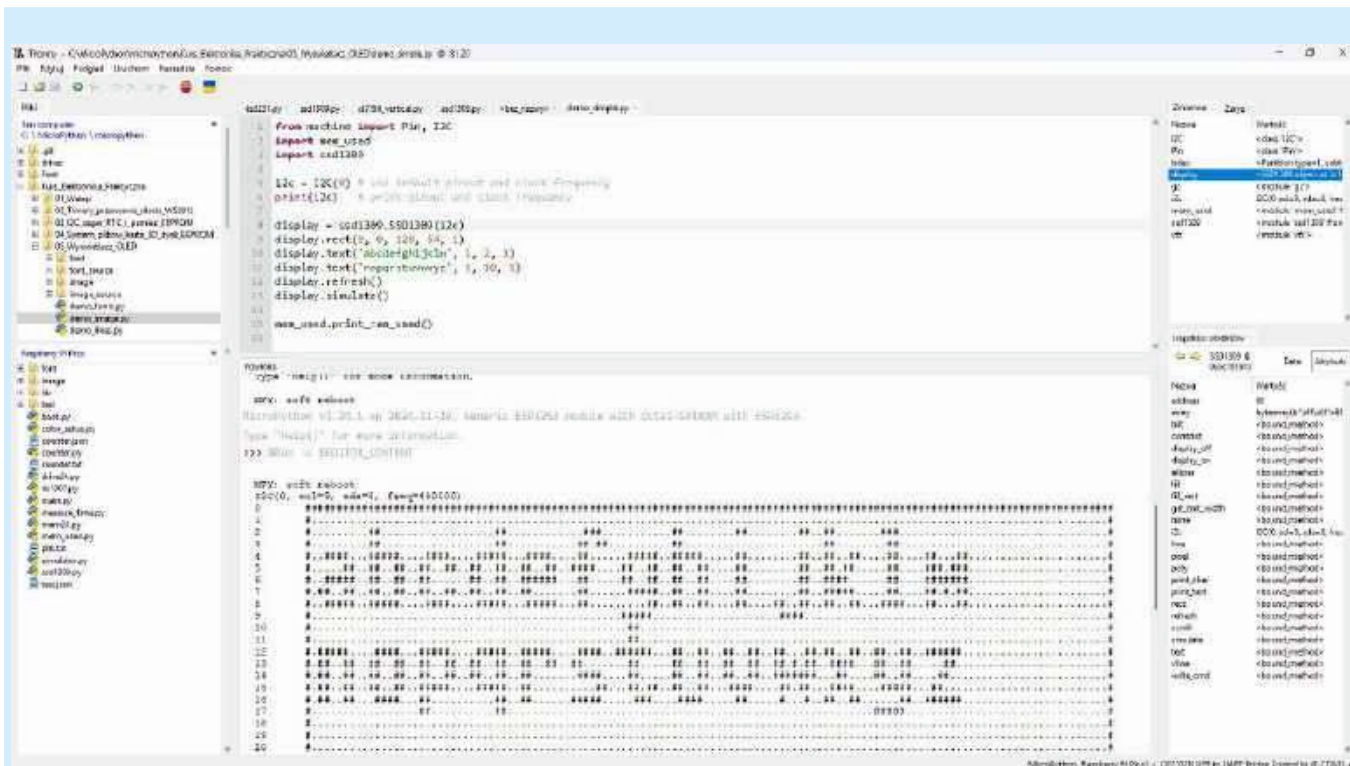
Program jest bardzo prosty – ma on wygenerować litery alfa-
betu korzystając z wbudowanego fontu w MicroPythona, a także nary-
sować prostokąt. Tak przygotowaną grafikę prześlemy do wyświetla-
cza oraz zasymulujemy na konsoli. Prześledźmy ten kod linia po linii.

Zaczynamy od zaimportowania klas **Pin** oraz **I2C** z modułu **ma-
chine** (linia 1), a także modułu **ssd1309** (linia 2), który omawiali-
śmy wcześniej. W linii 3 tworzymy instancję klasy **I2C**, korzystając
z interfejsu o numerze 0 i domyślnych pinów GPIO. Jeżeli chcesz
zastosować inne piny niż domyślne, podaj ich numery poprzez ar-
gumenty nazwane **sda** i **scl**. Linia 4 wyświetla na konsoli konfi-
gurację I2C – między innymi są to numery pinów używanych jako
SDA i SCL, a także częstotliwość zegara interfejsu.

W linii 5 tworzymy instancję klasy **SSD1309** z modułu **ssd1309**
i zapisujemy ją do zmiennej **display**. Poprzez argument przekazujemy



Rysunek 2. Układ współrzędnych wyświetlacza OLED ze sterownikiem SSD1309



Rysunek 3. Symulowany obraz na konsoli



Fotografia 3. Obraz na wyświetlaczu

do niej klasę `i2c`, która ma być wykorzystywana do komunikacji. W kolejnych liniach widzimy zakomentowane przykłady pokazujące, jak możemy zainicjalizować klasę, by wyświetlała obraz obrócony o 180° oraz jak zmienić adres na magistrali I²C.

Linia 6 to wywołanie metody `rect` z klasy `SSD1309` zapisanej w zmiennej `display`. Dwa pierwsze argumenty wskazują współrzędne lewego górnego narożnika prostokąta, czyli (0, 0). Dwa kolejne to jego wymiary, tzn. szerokość 128 pikseli i wysokość 64 piksele. Ostatni argument to kolor. W tym przypadku możliwe są tylko dwie opcje, czyli 1, co oznacza jasne piksele lub 0 – piksele wygaszone.

W liniach 7 i 8 generujemy napisy przy pomocy metody `text`. Pierwszym argumentem jest napis do wyświetlenia. Dwa kolejne to współrzędne x i y lewego górnego rogu napisu. Ostatni to kolor, podobnie jak w przypadku prostokąta.

Aby przesłać zawartość bufora obrazu do wyświetlacza, musimy wywołać metodę `refresh` (linia 9), natomiast by ów bufor

```
# Plik ep_logo_128x40.py

import framebuffer
ep_logo_128x40 = framebuffer.FrameBuffer(
    bytearray(b'\x00\xf8...\xc0\x00'),
    128, 40, framebuffer.MONO_VLSB)

Listing 3. Kod pliku ep_logo_128x40.py
```

wyświetlić na konsoli, wywołujemy metodę `simulate` (linia 10). Warto wspomnieć, że metody `refresh` i `simulate` pochodzą z klasy `SSD1309`, którą napisaliśmy sami, a metody `rect` i `text` są dziedziczone do `SSD1309` z klasy `FrameBuffer`.

Efekt działania kodu z listingu 2 pokazano na **rysunku 3** i **fotografii 3**.

Grafiki

Zajmijmy się teraz wyświetlaniem obrazów. Obrazki zapisujemy w osobnych plikach, a każdy z nich powinien być umieszczony w folderze `image` tak, jak pokazuje to **rysunek 4**. Oczywiście nie jest to jakaś żelazna zasada i można je umieszczać w dowolnych miejscach, ale proponowane podejście ułatwia pisanie kodu oraz automatyczne konwertowanie obrazów, o czym będzie w dalszej części artykułu.

Jak już pewnie zauważyłeś, każdy plik obrazu zapisany jest jako kod w Pythonie. Otwórzmy jeden z tych plików, na przykład `ep_logo_128x40.py`, którego kod pokazano na **listingu 3**. Składa się on tylko z dwóch linijek. Pierwszą jest import modułu `framebuf`. Druga to utworzenie zmiennej o nazwie takiej samej, jak nazwa pliku – tutaj wpisujemy instancję klasy `FrameBuffer`. Inicjalizujemy ją buforem pamięci, w którym znajduje się właściwa bitmapa, a w kolejnych argumentach podajemy szerokość, wysokość oraz format obrazu. To wszystko!

Zobaczmy teraz, w jaki sposób możemy zastosować tak przygotowane pliki obrazów. Przykład takich operacji pokazano



Rysunek 4. Katalog z obrazkami

```
# Plik demo_image.py

from machine import Pin, I2C
import ssd1309

from image.back_32x32 import *      # 1
from image.book_32x32 import *
from image.cancel_32x32 import *
from image.clock_32x32 import *
from image.down_32x32 import *
from image.ep_logo_128x40 import *
from image.hand_32x32 import *
from image.light_32x32 import *
from image.ok_32x32 import *
from image.settings_32x32 import *
from image.square_8x16 import *
from image.up_32x32 import *
from image.world_128x64 import *

i2c = I2C(0)
display = ssd1309.SSD1309(i2c)

display.blit(ok_32x32, 0, 0) # 2
display.blit(back_32x32, 0, 32)
display.blit(clock_32x32, 32, 0)
display.blit(settings_32x32, 32, 32)
display.blit(book_32x32, 64, 0)
display.blit(light_32x32, 64, 32)
display.blit(up_32x32, 96, 0)
display.blit(down_32x32, 96, 32)

# display.blit(world_128x64, 0, 0) # 3
# display.blit(world_128x64, 0, 0, 0) # 4

# display.blit(ep_logo_128x40, 0, 12)

display.refresh() # 5
```

Listing 4. Kod pliku demo_image.py

na listingu 4. Rozpoczynamy od zaimportowania klas **I2C** oraz **Pin** z modułu **machine** oraz modułu wyświetlacza **ssd1309**. Następnie importujemy obrazki (linia 1). Proponuję w tym przypadku użyć nieco innej składni niż zwykle, czyli **from katalog.plik import ***. Ów sposób wydaje mi się wygodniejszy, bo zaimportowane zmienne będą widoczne dokładnie tak, jakby były zmiennymi utworzonymi w pliku je importującym. Następnie tworzymy instancję klasy **I2C** oraz **SSD1309** – dokładnie tak samo, jak w poprzednich przykładach.

W linii 2 i kolejnych dodajemy obrazki do bufora wyświetlacza. Wykorzystujemy tutaj metodę **blit**, która pochodzi z klasy **FrameBuffer** i została odziedziczona do naszej klasy **SSD1309**. Przy pomocy trzech argumentów podajemy obiekt typu **FrameBuffer**, który przechowuje bitmapę do wyświetlenia oraz współrzędne x i y lewego górnego narożnika bitmapy.

W naszym przykładzie wyświetlamy osiem różnych ikon o rozmiarach 32x32 piksele, co powinno dać efekt widoczny na **fotografii 4**. W liniach 3 i 4 mamy polecenia wyświetlające kolejną bitmapę, reprezentującą mapę świata. Linia 4 różni się od 3 tylko tym, że ma dodatkowy argument o wartości 0. Powoduje to, że kolor czarny obrazu traktowany jest jako przezroczystość. W ten sposób można tworzyć warstwy z różnymi nakładającymi się na siebie grafikami.

Pozostaje już tylko linia 5, w której wywołujemy metodę **refresh**, aby przesłać bufor obrazu do wyświetlacza. Efekty uzyskane przy pomocy tego kodu pokazują **fotografie 5...7**.



Fotografia 4. Obraz z ośmioma ikonami



Fotografia 5. Mapa świata

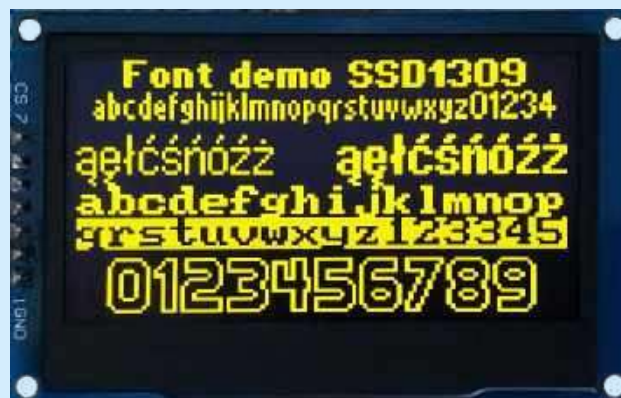


Fotografia 6. Logo „Elektroniki Praktycznej”

Jak wygenerować pliki z obrazkami? Niestety autorzy MicroPythona nie dostarczają żadnych narzędzi do tego celu i trzeba je zrobić samemu. Zaprezentuję tutaj moje rozwiązanie, które opracowałem na własne potrzeby jakiś czas temu.

W katalogu naszego projektu musimy utworzyć dwa podkatalogi. W pierwszym z nich, o nazwie **image_source**, umieszczać będziemy wszystkie potrzebne grafiki w formacie BMP, zapisane jako bitmapy 16-kolorowe. W tym celu możemy posłużyć się programem Paint, w którym format obrazu i paletę kolorów wybiera się podczas zapisywania pliku. W drugim katalogu, o nazwie **image**, znajdą się pliki po konwersji.

Właściwą konwersję wykonuje skrypt **_convert_images.py** (w moich projektach często stosuję znak „_” na początku nazwy pliku, aby wyróżnić pliki Pythona przeznaczone do wykonania na komputerze, a nie ESP32), który umieścić musimy w katalogu głównym projektu. Jego działanie jest bardzo proste – otwiera po kolei wszystkie pliki *.bmp z katalogu **image_source**, przekształca je i wynik zapisuje w katalogu **images**, w plikach o analogicznych nazwach, ale z rozszerzeniem *.py.



Fotografia 7. Obraz uzyskany przez kod z listingu 6

koktajl niusów



Firma Creotech Instruments zbuduje narodową konstelację satelitarną CAMILA

Creotech Instruments podpisała umowę na realizację narodowej konstelacji satelitarnej CAMILA (Country Awareness Mission in Land Analysis). Projekt obejmuje budowę trzech satelitów: radarowego, optycznego o wysokiej rozdzielczości oraz optycznego o niższej rozdzielczości. Umowa przewiduje również opcje rozszerzające zakres prac, w tym dostawę czwartego satelity oraz zapewnienie obsługi konstelacji na orbicie; opcje te mogą podnieść wartość kontraktu i zwiększyć ogólną atrakcyjność przedsięwzięcia.

Celem programu jest wzmocnienie bezpieczeństwa narodowego oraz zwiększenie niezależności Polski na arenie międzynarodowej. Założono możliwie najszersze wykorzystanie komponentów krajowej produkcji: rozwiązanie wyłonione w otwartym przetargu bazuje co najmniej w 90 procentach na podsystemach stworzonych w Polsce, a prawa własności intelektualnej do tych rozwiązań pozostają w kraju. Projekt zakłada nie tylko uruchomienie operacyjnego systemu składającego się z minimum trzech satelitów, lecz także konsolidację i wsparcie polskiej myśli technicznej.

Wiele podsystemów planowanych do wdrożenia w ramach projektu CAMILA przeszło już weryfikację w przestrzeni kosmicznej podczas misji EagleEye zrealizowanej przez Creotech Instruments w 2024 roku. Na orbitę trafił wówczas satelita sprzętowo oparty na polskiej platformie HyperSat. Przeprowadzone testy pozwoliły zidentyfikować i usunąć błędy konstrukcyjne, co ma ograniczyć ryzyko i ułatwić udany debiut systemu CAMILA.

<https://tiny.pl/rdr-kytb>

Firmy Schneider Electric i NVIDIA przyspieszają rozwój fabryk AI

Schneider Electric i NVIDIA ogłosiły rozszerzenie partnerstwa oraz prezentację nowej linii rozwiązań infrastrukturalnych dla centrów danych projektowanych pod zastosowania sztucznej inteligencji. Wśród zaprezentowanych nowości znalazły się najnowsze wersje EcoStruxure Pod oraz Rack Infrastructure, czyli skalowalne moduły mające skracać czas wdrożenia centrów danych przygotowanych do pracy w modelu AI-ready.

Jednym z kluczowych elementów oferty jest system szaf serwerowych opracowany z myślą o platformie NVIDIA GB200 NVL72, opartej na modularnej architekturze NVIDIA MGX. Wraz z tym



rozwiązaniem Schneider Electric po raz pierwszy łączy do ekosystemów NVIDIA HGX oraz MGX, wspierając rozwój infrastruktury dla zaawansowanych zastosowań AI. Ogłoszone rozwiązania stanowią kontynuację współpracy obu firm, która w 2025 roku zaowocowała powstaniem pierwszego na świecie cyfrowego bliźniaka systemów elektroenergetycznych w fabrykach AI.

Zaprezentowano także projekty referencyjne obejmujące obszary zasilania oraz chłodzenia cieczą; część z tych rozwiązań wykorzystuje technologię Motivair, producenta dedykowanych systemów, który dołączył do Schneider Electric w marcu 2025 roku. Strony zapowiadają dalsze inicjatywy w ramach partnerstwa, w tym rozwój kolejnych produktów i projektów referencyjnych oraz wsparcie dla budowy tzw. fabryk AI, z wykorzystaniem ogłoszonych modułów infrastrukturalnych i kompatybilności z platformami NVIDIA.

<https://tiny.pl/r25xh4sq>

Firma Emitel testuje standard telewizji mobilnej 5G Broadcast

Emitel prowadzi testy nowego standardu telewizji mobilnej 5G Broadcast, utrzymując emisję sygnału programów TV Puls i PULS 2. Transmisja realizowana jest z jednego z obiektów nadawczych firmy i obejmuje obszar aglomeracji warszawskiej. Celem przedsięwzięcia jest ocena możliwości technologicznych oraz potencjału wdrożeniowego tego rozwiązania. Jest to pierwszy w Polsce projekt, w którym do dystrybucji sygnału telewizyjnego wykorzystano istniejącą infrastrukturę nadawczą, bez konieczności logowania do internetu czy stosowania kart SIM.



Technologia 5G Broadcast może znaleźć zastosowanie w wielu obszarach, od darmowej telewizji mobilnej, przez radio cyfrowe, po systemy powiadamiania kryzysowego. Może być także integrowana z usługami interaktywnymi dzięki modelowi hybrydowemu, umożliwiającemu udostępnianie dodatkowych treści na życzenie. Według ekspertów w nadchodzących latach 5G Broadcast stanie się ważnym narzędziem dla nadawców, operatorów telekomunikacyjnych i platform streamingowych. Pozwoli na jednoczesną dystrybucję treści do dużej liczby smartfonów, tabletów i odbiorników samochodowych, bez konieczności zestawiania indywidualnych połączeń. Rozwiązanie zapewnia niezawodny odbiór także przy dużym obciążeniu sieci lub jej czasowej niedostępności, co czyni je szczególnie przydatnym w przypadku masowych transmisji, wydarzeń na żywo oraz w sytuacjach nadzwyczajnych.

Technologia 5G Broadcast może znaleźć zastosowanie w wielu obszarach, od darmowej telewizji mobilnej, przez radio cyfrowe, po systemy powiadamiania kryzysowego. Może być także integrowana z usługami interaktywnymi dzięki modelowi hybrydowemu, umożliwiającemu udostępnianie dodatkowych treści na życzenie. Według ekspertów w nadchodzących latach 5G Broadcast stanie się ważnym narzędziem dla nadawców, operatorów telekomunikacyjnych i platform streamingowych. Pozwoli na jednoczesną dystrybucję treści do dużej liczby smartfonów, tabletów i odbiorników samochodowych, bez konieczności zestawiania indywidualnych połączeń. Rozwiązanie zapewnia niezawodny odbiór także przy dużym obciążeniu sieci lub jej czasowej niedostępności, co czyni je szczególnie przydatnym w przypadku masowych transmisji, wydarzeń na żywo oraz w sytuacjach nadzwyczajnych.

<https://tiny.pl/c1gsrz6k>

Niezastąpiona karta graficzna Condor GR6S-AD2000 od EIZO

Firma EIZO wprowadziła do swojej oferty kartę graficzną Condor GR6S-AD2000, zaprojektowaną z myślą o przechwytywaniu wideo w standardzie SOSA oraz obsłudze wysokowydajnych systemów obliczeniowych wymagających intensywnego przetwarzania danych przy użyciu sztucznej inteligencji. Urządzenie obsługuje do sześciu stru-

mieni wideo 12G-SDI, a jego elastyczna architektura umożliwia pobieranie, przetwarzanie i transmisję obrazu.

Sercem karty jest procesor graficzny NVIDIA RTX 2000 Ada Generation wyposażony w 8 GB pamięci GDDR6, 3072 rdzenie CUDA, 96 rdzeni Tensor oraz 24 rdzenie RT. Dzięki obsłudze kodowania AV1 i zaawansowanym możliwościom przechwytywania wideo, karta zapewnia efektywną kompresję o niskich opóźnieniach oraz umożliwia strumieniowanie materiału wideo, co jest szczególnie istotne w aplikacjach C5ISR. Zastosowane interfejsy PCI Express Gen 4 (x4, x8 lub x16) gwarantują wysoką przepustowość, a wsparcie dla bifurkacji i łańcuchowania pozwala na elastyczne konfigurowanie systemu w zależności od potrzeb.

Condor GR6S-AD2000 spełnia wymagania standardu SOSA, wspierając profile modułów 14.6.11-0 i 14.6.13-0, a ponadto jest zgodna ze specyfikacjami środowiskowymi VITA 47.1 dotyczącymi odporności na temperaturę, wstrząsy i wibracje. Dzięki temu karta sprawdzi się w wymagających warunkach pracy, zapewniając jednocześnie stabilność i wysoką wydajność w zadaniach związanych z nieprzerwanym przetwarzaniem obrazu.

<https://tiny.pl/fps11p33>



Adapter sieciowy Linksys MAX-STREAM AC600 z technologią MU-MIMO

Linksys wprowadza adapter MAX-STREAM AC600, przeznaczony do komputerów stacjonarnych i mobilnych z systemem Windows. Urządzenie ma umożliwić modernizację maszyn wyposażonych w starsze karty Wi-Fi (np. wyłącznie Wireless-N) poprzez obsługę standardu Wireless-AC i technologii Next-Gen MU-MIMO. Producent deklaruje, że dedykowane strumienie Wi-Fi pomagają utrzymać płynność transmisji podczas oglądania materiałów wideo w jakości HD, rozgrywek online oraz jednoczesnego pobierania wielu plików. Kompaktowa obudowa sprawia, że adapter jest dyskretnym dodatkiem do laptopa.

MAX-STREAM AC600 pracuje w pasmach 2,4 GHz i 5 GHz; w przypadku pierwszego z nich przepustowość wynosi 150 Mb/s, a drugiego: 433 Mb/s. Urządzenie może działać na obu pasmach w celu ograniczania zakłóceń, z możliwością przypisania nowszych urządzeń Wireless-AC do 5 GHz, a starszych urządzeń Wireless-N i -G do 2,4 GHz. Zastosowana technologia Beamforming ukierunkowuje sygnał na klienta, co ma zwiększać efektywność komunikacji bezprzewodowej. Adapter współpracuje z szeroką gamą routerów, przy czym maksymalną wydajność zapewnia zestawienie go z routerem MU-MIMO.

Instalacja przebiega przez podłączenie portu USB 2.0 lub 3.0; sterownik jest automatycznie pobierany z serwerów Microsoft, a w razie potrzeby można skorzystać z dołączonej płyty CD. Wymagany jest komputer z napędem optycznym CD lub DVD, systemem Windows w wersji 7 lub nowszej oraz aktywne łącze internetowe. Adapter nie wymaga logowania do dodatkowych usług producenta i działa jako standardowe urządzenie klienckie Wi-Fi.

https://linksys.pl/?page_id=1341



Drukarka SureColor SC-S8100 dla profesjonalistów z branży etykiet

Firma Epson wprowadziła do swojej oferty model SureColor SC-S8100, który dołącza do serii SC-S7100 i SC-S9100, tworząc pełną



linię rozwiązań przeznaczonych dla profesjonalnego rynku druku etykiet. Nowe urządzenie stanowi rekomendowanego następcę popularnej serii SC-S60600 i wprowadza szereg udoskonaleń konstrukcyjnych oraz technologicznych, których brak w starszych modelach. Drukarka wykorzystuje zestaw atramentów UltraChrome GS3 w konfiguracji CMYK z dodatkiem Light Cyan i Light Magenta, co pozwala uzyskać szeroką gamę barw oraz zminimalizować efekt bandingu przy zachowaniu wysokiej jakości obrazu.

Centralnym elementem urządzenia jest głowica PrecisionCore Micro TFP z 9 600 dyszami, oferująca o 30% wyższą wydajność w porównaniu z poprzednimi modelami. Konstrukcja drukarki została obniżona do wysokości 1,021 m, a przezroczysta obudowa z niskim profilem ułatwia obserwację procesu druku i kontrolę zużycia materiałów. Obsługę wspiera panel dotykowy o przekątnej 4,3 cala. Wydajność pracy zwiększa technologia weryfikacji dysz (NVT), która automatycznie monitoruje ich stan i koryguje nieprawidłowości, a wbudowany system czyszczący chroni podzespoły przed kurzem i zanieczyszczeniami.

SureColor SC-S8100 jest w pełni kompatybilna z platformą Epson Cloud Solution PORT, umożliwiającą zdalne monitorowanie efektywności drukarki, zużycia materiałów i statusów serwisowych. Rozwiązanie to pozwala na bieżącą kontrolę procesu druku i optymalizację kosztów eksploatacyjnych, co ma znaczenie dla firm działających w dynamicznie zmieniających się warunkach rynkowych.

<https://tiny.pl/sw9g8zdm>

Inteligentna kamera zewnętrzna Smart EYE 400

Kamera zewnętrzna Smart EYE 400 została zaprojektowana z myślą o zapewnieniu skutecznego monitoringu w różnych warunkach środowiskowych. Wyposażono ją w obiektyw współpracujący z technologią IR CUT, co pozwala uzyskiwać wyraźny obraz zarówno w dzień, jak i w nocy. Konstrukcja urządzenia jest odporna na wilgoć i działanie wysokich temperatur, co umożliwia pracę w wymagającym otoczeniu.

Model Smart EYE 400 oferuje funkcję wykrywania ruchu, która pozwala na bieżąco rejestrować zmiany w monitorowanym obszarze. Kamera może być montowana na ścianie lub suficie, zapewniając precyzyjne wykrywanie w odległości od 1,5 do 15 metrów przy kącie nachylenia od 30° do 60°.

Połączenie parametrów optycznych, detekcji ruchu oraz odporności na czynniki zewnętrzne sprawia, że Smart EYE 400 znajduje zastosowanie w systemach monitoringu, w których priorytetem jest niezawodna ochrona oraz stała kontrola otoczenia.

<https://tiny.pl/8v7qs92y>





Najnowszy zwijacz kabla mikrofonowego e-spool od firmy igus do zastosowań scenicznych i eventowych

Firma igus wprowadziła zwijacz mikrofonu e-spool z dedykowanym przewodem XLR, zaprojektowany w celu ograniczenia ryzyka awarii podczas kluczowych wydarzeń muzycznych czy audiowizualnych. Konstrukcja mechaniczna eliminuje działanie sił rozciągających na żyły przewodu, a wbudowane elementy zwijania wyposażono w funkcję automatycznego wyłączenia, która zatrzymuje ruch w przypadku nietypowego obciążenia, na przykład wskutek przypadkowego zaczepienia. Urządzenie jest zasilane elektrycznie i umożliwia automatyczne prowadzenie kabli mikrofonowych po suficie do długości 30 metrów, uruchamiane jednym przyciskiem.

Kluczowym elementem rozwiązania jest tor prowadzenia przewodu: kabel wyprowadzany z boku bębna trafia do połączonych ogniw przewodnika wykonanych z polimerów o wysokiej wytrzymałości. Podczas obrotu bębna ogniw zazębiają się i wykonują spiralny ruch, co stabilizuje ułożenie przewodu i ma sprzyjać jakości sygnału oraz niezawodności transmisji. Zastosowano dedykowany silnik elektryczny do precyzyjnego opuszczania przewodów w kierunku podłogi na całej długości do 30 metrów. Wersja stereo umożliwia równoczesne prowadzenie dwóch przewodów mikrofonowych.

Zwijacz e-spool może być obsługiwany za pomocą sterownika przewodowego albo integrowany z nadrzędnymi systemami sterowania poprzez sieć WLAN lub inne interfejsy radiowe. Zestaw funkcji obejmujących mechaniczne odciążenie przewodu, automatyczne zabezpieczenie przed przeciążeniem oraz sterowanie elektryczne to odpowiedź na potrzeby instalatorów, którzy poszukują sposobu na bezpieczne i uporządkowane prowadzenie kabli mikrofonowych w przestrzeni sufitowej.

<https://tiny.pl/7b340f3p>



Nowa klasa wysokoefektywnych materiałów emitujących światło opracowana przez polsko-brytyjski zespół

Zespół naukowców z Instytutu Chemii Fizycznej Polskiej Akademii Nauk i Wydziału Chemicznego Politechniki Warszawskiej, we współpracy z badaczami z Uniwersytetu Cambridge, zsyntetyzował

nową klasę wysoce luminescencyjnych kompleksów glinoorganicznych. Czerpiąc z wcześniejszych wyników i materiałów odniesienia, opracowano serię niespotykanych dotąd tetramerycznych kompleksów alkilowoglinowych zawierających ligandy wywodzące się z aminokwasów aromatycznych. Badania fotofizyczne wykazały, że antranilany oparte na glinie charakteryzują się bardzo wysoką wydajnością kwantową fotoluminescencji w stanie stałym; autorzy podkreślają, że może ona należeć do najwyższych raportowanych dla tego typu układów.

Obliczenia kwantowo-chemiczne dostarczyły wglądu w naturę przejść elektronowych i pozwoliły zidentyfikować fragmenty cząsteczek, które w największym stopniu odpowiadają za obserwowane właściwości fotofizyczne. Wykazano, że subtelne modyfikacje ligandów podnoszą sprawność emisji poprzez tłumienie niepożądanych ścieżek relaksacji. W ciele stałym istotną rolę odgrywają niekowalencyjne oddziaływania wewnątrz- i międzycząsteczkowe, które pomagają utrzymać integralność strukturalną w stanie wzbudzonej i ograniczają zniekształcenia prowadzące do wygaszania fluorescencji. Odpowiednio kontrolowana agregacja cząsteczek zwiększa sztywność układu, co sprzyja uzyskaniu wysokiej luminescencji.

Autorzy oceniają, że opracowane kompleksy stanowią ważny krok w projektowaniu nowych, łatwo dostępnych i wydajnych materiałów fluorescencyjnych. Prostota modyfikacji szkieletu ligandów stwarza możliwość dalszego zwiększania stabilności chemicznej oraz precyzyjnego dostrajania właściwości optycznych. Wyniki przybliżają zastosowania praktyczne w technologiach optoelektronicznych, w tym w diodach OLED, ekranach wyświetlaczy oraz czujnikach.

<https://tiny.pl/vrvv-kb9x>

Wielopunktowy rejestrator wilgotności HUMi1REG2.0 do monitoringu parametrów środowiskowych

HUMi1REG2.0 to wielopunktowy rejestrator przeznaczony do stałego monitorowania parametrów środowiskowych oraz przeglądania danych historycznych. Bieżące wartości uśredniane i zapisywane do wbudowanej pamięci. Urządzenie standardowo rejestruje uśrednione wartości z pełnych godzin; po wypełnieniu pamięci (ok. 1,5 roku nieprzerwanej rejestracji) najstarsze rekordy są automatycznie nadpisywane. Dostęp do archiwów możliwy jest w formie tabel i wykresów, a dane można eksportować na karty SD w formacie CSV. Rejestrator udostępnia wyjście przekaźnikowe do sygnalizacji przekroczeń zadanych progów, a wykryte problemy z czujnikami raportuje natychmiast.

Wbudowany moduł Wi-Fi zapewnia łączność i automatyczne przesyłanie danych do wybranych usług chmurowych, a aktualne i archiwalne odczyty można przeglądać w aplikacji mobilnej NtroMonit. Do lokalnej obsługi i komunikacji służy ekran dotykowy o przekątnej 5 cali. HUMi1REG2.0 jest przeznaczony do zastosowań wymagających utrzymania ściśle określonych wartości wilgotności lub temperatury, m.in. podczas profilowania temperatury pieców.

<https://tiny.pl/f6bs1dx6>

Jakub Tyburski
jakub.tyburski@elportal.pl



Temat Numeru: Maszyny i wyposażenie do produkcji

Nowoczesna produkcja elektroniki to zaawansowany i wieloetapowy proces, w którym dokładność, powtarzalność i skrupulatna kontrola jakości mają fundamentalne znaczenie dla jakości finalnego produktu. Od nakładania pasty lutowniczej aż po końcową inspekcję – każdy etap wymaga odpowiedniego parku maszynowego i właściwego przygotowania stanowisk pracy. W artykule prezentujemy najważniejsze urządzenia wykorzystywane w montażu powierzchniowym SMT: automaty pick&place, piece rozplływowe (reflow), drukarki szablone do pasty lutowniczej czy też systemy inspekcji optycznej (AOI) i rentgenowskiej (X-Ray). Omawiamy również rolę robotów lutowniczych w procesach THT oraz znaczenie wyposażenia antystatycznego stref EPA, w tym mebli ESD, mat, opasek, a także organizacji stanowisk zgodnych z normami ochrony przed wyładowaniami elektrostatycznymi. Artykuł będzie cennym źródłem wiedzy dla osób planujących uruchomienie lub modernizację linii produkcyjnej.



Elektronika w Praktyce: Transformatory do różnych aplikacji

Pomimo upływu lat, w niektórych zastosowaniach transformatory wciąż są nie do zastąpienia. Czasem możemy pójść na skróty i zrezygnować z nich w niektórych typowych aplikacjach, tym bardziej że współczesna elektronika nieco spycha te elementy na margines. Przetwornice impulsowe stają się już na tyle doskonałe, że transformatory w roli konwerterów napięcia często nie są już potrzebne, choć nie zawsze da się je tak łatwo zastąpić. Ponadto miniaturyzacja nie może objąć transformatorów w tak znaczącym stopniu, jak ma to miejsce w przypadku innych gałęzi elektroniki, zwłaszcza półprzewodników. Które z podzespołów są transformatorami, choć tak nie wyglądają? Jak zrobić transformator bez drutu? I dlaczego, mimo ich ograniczeń, powinniśmy je doceniać? Zapraszamy do lektury październikowej odsłony „Elektroniki w Praktyce”!



Kompaktowy licznik czasu pracy z detekcją przepływu prądu

Czas pracy maszyny, mierzony podczas pobierania przez nią energii elektrycznej, to jeden z podstawowych parametrów warunkujących prawidłową pracę urządzeń wymagających okresowego serwisu. Prezentowany układ umożliwia łatwą realizację takiego pomiaru, a dodatkowo ma niewielkie wymiary. Urządzenie mierzy czas pracy maszyny w okresach zwiększonego poboru prądu. Rozdzielczość pomiaru wynosi 1 s, choć wynik jest prezentowany w jednej z trzech bardziej czytelnych form: dni, dni i godzin lub godzin i minut, w zależności od wymagań. Pomiar prądu jest realizowany przez przekładnik prądowy, więc nie ma galwanicznego połączenia licznika z siecią.



Najważniejsze parametry:

- maksymalna pojemność licznika: 9999 dni lub 99 dni i 23 godziny lub 99 godzin i 59 minut,
- pomiar pobieranego prądu poprzez izolujący galwanicznie przekładnik prądowy,
- możliwość ustawienia progu nieczułości urządzenia w zakresie 240 mA...24 A wartości skutecznej pobieranego prądu,
- sygnalizacja przepełnienia licznika,
- czytelny, czterocyfrowy wyświetlacz 7-segmentowy LED,
- możliwość wyzerowania wskazań,
- zapamiętywanie zmierzonego czasu podczas zaniku zasilania,
- pobór prądu około 20 mA,
- zasilanie napięciem stałym 9...30 V.

Wykaz firm ogłaszających się w tym numerze „Elektroniki Praktycznej”

AKSOTRONIK	13
AVT	5, 49, 55, 71
BORNICO	51
COMPUTER CONTROLS	7
CONRAD ELECTRONIC	21, 88
FAULHABER	15, 35
FERYSER	31
LASTENIC LASER	11
PHOENIX CONTACT	22, 23
SEMICON	9

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 Guguła, tel. 22 257 84 64

Szef Pracowni Konstrukcyjnej:
Jakub Sobański

Zespół marketingu i reklamy:
Katarzyna Guguła, Bożena Krzykawska,
Grzegorz Krzykowski

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 Tatuś, Jakub Tyburski, Robert Wołgajew

Uwaga!
Kontakt z wymienionymi osobami jest możliwy via e-mail,
według schematu: imię.nazwisko@ep.com.pl

DTP, 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

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”.





Tak! Pełna moc zamiast przerwy w dostawie prądu. Z Conrad.

Wysokiej jakości technologia pomiarowa i odpowiednie części zamienne



conrad.pl/awarie-elektryczne

All parts of success

CONRAD

eprasa.pl 2314d257b7