

ELEKTRONIKA PRAKTYCZNA

● MIĘDZYNARODOWY MAGAZYN ELEKTRONIKÓW KONSTRUKTORÓW ● LIPIEC/SIERPIEŃ ● 7-8/2026 ●



wybierz właściwy OSCYLOSKOP

Z praktyki ekspertów

- Idealny wzmacniacz operacyjny, który oscyluje
- Technika wielkiej częstotliwości: Mit 50 omów
- Ognia i systemy BMS:
Stan naładowania to nie napięcie
- θ_{JA} z noty katalogowej to fikcja
- Profesor Christophe Basso: Stabilizacja pętli trójfazowego prostownika Vienna, część 1

W pracowni

- Radiodbiorniki FM
- Czujniki pojemnościowe położenia
- Sterowniki wentylatorów. Dwuwentylatorowy sterownik PWM z izolowaną przetwornicą flyback 6 W
- Zgrzewarki oporowe

Repetitorium

- Oscyloskopy 2026, część 3
- Astabilne układy czasowe

Kity AVT

- Miniaturowy wzmacniacz audio w klasie D (AVT6099)

Moduły z Chin

- Luckfox Pico 86 Panel – moduł do automatyki domowej z AI, montowany w standardowej puszcze elektrycznej

AI w elektronice

- TinyML: duże uczenie w małym układzie

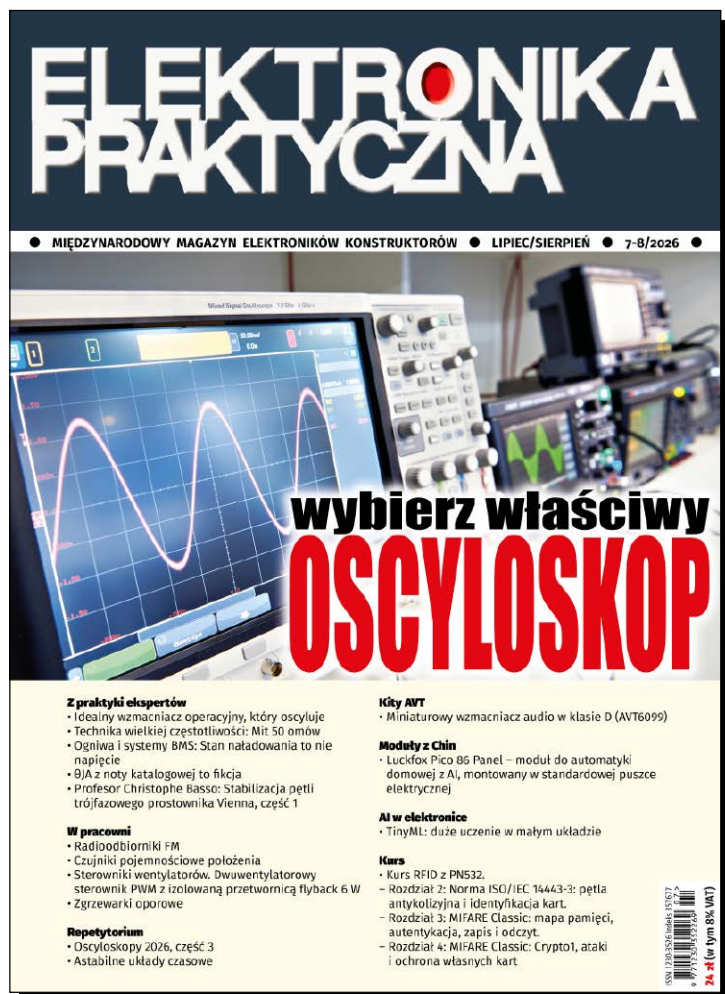
Kurs

- Kurs RFID z PN532.
- Rozdział 2: Norma ISO/IEC 14443-3: pętla antykolizyjna i identyfikacja kart
- Rozdział 3: MIFARE Classic: mapa pamięci, autentykacja, zapis i odczyt
- Rozdział 4: MIFARE Classic: Crypto1, ataki i ochrona własnych kart

E-prenumerata Elektroniki Praktycznej

Twoja wiedza zawsze pod ręką!

Zyskaj 30% rabatu na roczny dostęp cyfrowy (PDF)



W e-prenumeracie
zapłacisz tylko:

~~240,00 zł~~

168,00 zł/rok

10 wydań w roku
16,80 zł za numer



Zamów prenumeratę na
www.UlubionyKiosk.pl
lub zeskanuj kod QR
i zaprenumeruj w 1 minutę!

Dlaczego warto?

- ✓ **taniej niż w sprzedaży pojedynczej:** rok dostępu za 168 zł zamiast 240 zł
– wychodzi 16,80 zł za numer i 72 zł oszczędności w skali roku
- ✓ **nowy numer automatycznie w dniu premiery:** trafia na Twoje konto bez pilnowania kolejnych wydań i bez ryzyka, że któryś przegapisz
- ✓ **komplet bez luk:** kursy, repetytoria i cykle „z praktyki ekspertów” prowadzą krok po kroku przez kilka numerów – prenumerata gwarantuje ciągłość całej serii
- ✓ **cała biblioteka pod ręką:** wszystkie numery na tablecie, laptopie i smartfonie, w przeszukiwalnym pliku PDF
- ✓ **przywileje prenumeratora:** specjalne zniżki na pozostałe tytuły z oferty serwisu UlubionyKiosk.pl

AVTKorporacja sp. z o.o., ul. Leszczynowa 11, 03-197 Warszawa
prenumerata@avt.pl | 22 257 84 22 (godz. 10.00–14.00)
rachunek bankowy: ING Bank Śląski 18 1050 1012 1000 0024 3173 1013

Oscyloskop to dziś inne zwierzę

Wyposażenie instrumentalne warsztatu elektronika zmienia się z upływem lat, ale żaden przyrząd nie zmienił się tak radykalnie jak oscyloskop. Elektronik wykształcony 20..30 lat temu, postawiony przed współczesnym stanowiskiem, poradzi sobie z zasilaczem, generatorem, multimetrem – to wciąż te same przyrządy, z tą samą logiką obsługi. Oscyloskop sprawi mu jednak spory kłopot, bo to dziś zupełnie inna konstrukcja i inna filozofia pracy.

Dawny oscyloskop analogowy był w istocie oscylografem: wiązka elektronów odchylana w czasie rzeczywistym przez mierzony sygnał, obserwowana wprost na ekranie lampy. Nastawiało się pasmo, podstawę czasu i poziom wyzwalania – i widziało się sygnał taki, jaki jest. Współczesny DSO tego sygnału nie pokazuje wprost: próbkuje go przetwornikiem analogowo-cyfrowym, zapisuje w pamięci akwizycji i odtwarza na ekranie. Nic z tego, co widać, nie jest już bezpośrednio sygnałem – jest wynikiem obliczeń.

Z tą zmianą przyszedł cały słownik, którego dawne stanowisko nie znało: częstotliwość próbkowania obok pasma, aliasing i interpolacja, głębokość pamięci akwizycji, tryby hi-res, peak-detect i uśredniania, cyfrowe wyzwalanie zaawansowane, pamięć segmentowa, rozdzielczość 12 bitów i ENOB, kanały logiczne (MSO), dekodowanie protokołów, FFT i matematyka przebiegów. Pokręta w dużej mierze ustąpiły miejsca menu na ekranie dotykowym. Przyrząd stał się komputerem pomiarowym i wymaga innego modelu myślenia niż lampa z wiązką.

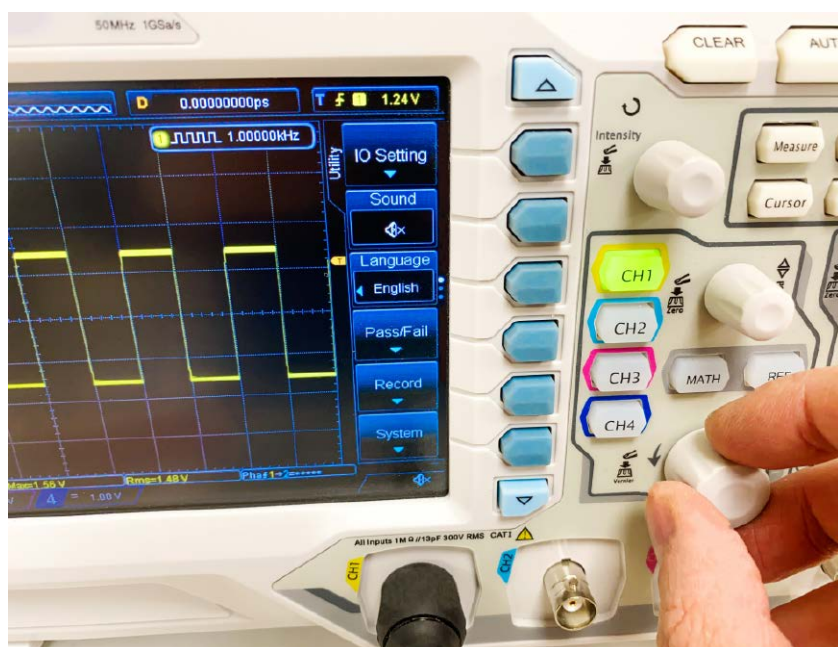
Najlepiej widać tę różnicę na wyzwalaniu. Na oscyloskopie analogowym był to jeden potencjometr poziomu; dziś to rozbudowane menu – wyzwalanie zboczem, szerokością impulsu, na warunku logicznym, na zdekodowanej ramce protokołu, wreszcie strefowe. Dawny elektronik odruchowo szuka „pokrętła triggera”, a znajduje drzewo opcji. Część nawyków wciąż jednak procentuje: kompensacja sondy, zależność czasu narastania od pasma czy pętla masy przy szybkich zboczach to fundamenty, które się nie zmieniły. Nowa jest przede wszystkim warstwa akwizycji i obróbki danych – i to ona wymaga dołożenia wiedzy.

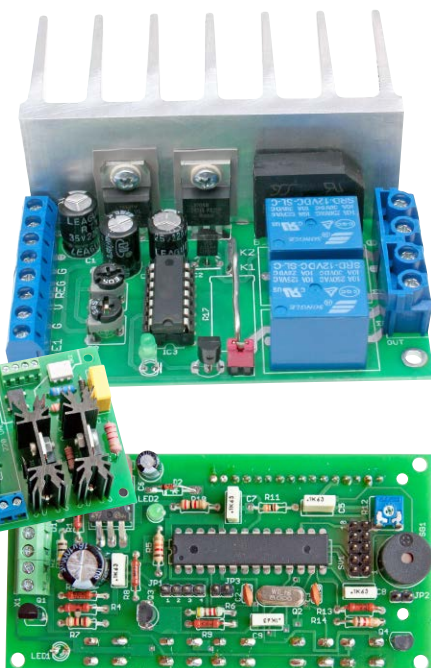
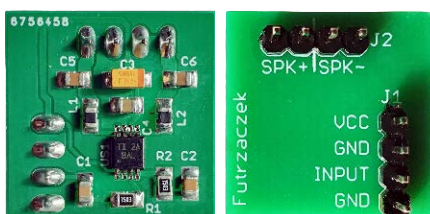
Ta zmiana to nie tylko nowa terminologia, lecz nowe możliwości pomiarowe. Głęboka pamięć i szybkie odświeżanie pozwalają wyłapać rzadki, jednorazowy glitch, który na lampie analogowej przemknąłby niezauważony. Wbudowane dekodery pokazują ramki I²C, SPI czy UART wprost nad przebiegiem elektrycznym, bez osobnego analizatora. Natywne 12 bitów ujawnia tętnienie na tle dużego napięcia stałego, którego 8-bitowy przyrząd nie rozdzielił od szumu. Pamięć segmentowa zapisuje tysiące krótkich zdarzeń, nie zapełniając się martwym czasem między nimi. To dawny oscyloskop robił albo gorzej, albo wcale – i dlatego nowego przyrządu trzeba się nauczyć od nowa, a nie tylko do niego przywyknąć.

Dlatego poświęciliśmy oscyloskopowi trzyczęściowe repetytorium, którego ostatnią część zamieszczamy w tym numerze. Prowadzi ono dokładnie po tej drodze – od lampy Brauna do 12 bitów na biurku konstruktora. Część pierwsza (EP 05/2026) to krótka historia oraz architektura współczesnego DSO: tor pionowy, przetwornik A/C i pamięć akwizycji. Część druga (EP 06/2026) schodzi na poziom próbkowania, wyzwalania i akwizycji oraz sond i integralności sygnału – tam, gdzie „1 GSa/s w nocy” nie zawsze oznacza 1 GSa/s na kanale. Część trzecia, w tym wydaniu, pokazuje, co zmieniło się w latach 2020...2026: wejście 12 bitów i kanałów logicznych do klasy budżetowej – i jak z tej wiedzy skorzystać przy zakupie.

Adresatem cyklu jest zarówno ten, kto uczył się oscyloskopu na lampie i chce nadrobić trzy dekady, jak i ten, kto zaczyna dziś i chce zrozumieć, za co właściwie płaci. Bo choć oscyloskop stał się innym zwierzęciem, jego rola się nie zmieniła: to wciąż przyrząd, bez którego elektroniki uprawiać się nie da.

Wiesław Marciniak





NIE PRZEOCZ

Nowe podzespoły 6

AI W ELEKTRONICE

AI w pracowni elektronika konstruktora, część 3
TinyML: duże uczenie w małym układzie 14

Z PRAKTYKI EKSPERTÓW

Idealny wzmacniacz operacyjny, który oscyluje 19
Technika wielkiej częstotliwości: Mit 50 omów 29
Ogniwa i systemy BMS: Stan naładowania to nie napięcie 40
 θ_{JA} z noty katalogowej to fikcja 48
Profesor Christophe Basso:
Stabilizacja pętli trójfazowego prostownika Vienna, część 1 55

PREZENTACJE

Nowoczesna technika pomiarowa
wspiera precyzyjną diagnostykę i serwis 66
Oscyloskop, który wchodzi w wysokie napięcie 90
Co potrafi nowoczesny oscyloskop.
Sześć funkcji, które w praktyce zmieniają pracę inżyniera 108

W PRACOWNI

Radioodbiorniki FM 68
Czujniki pojemnościowe położenia 72
Sterowniki wentylatorów. Dwuwentylatorowy sterownik PWM
z izolowaną przetwornicą flyback 6 W 75
Zgrzewarki oporowe 80

MODUŁY Z CHIN

Luckfox Pico 86 Panel – moduł do automatyki domowej z AI,
montowany w standardowej puszcze elektrycznej 85

KITY AVT

Miniaturowy wzmacniacz audio w klasie D 88

REPETYTORIUM

Oscyloskopy 2026 – od lampy Brauna
do 12 bitów na biurku konstruktora, część 3 97
Astabilne układy czasowe 112

KURS

Kurs RFID z PN532, rozdział 2.
Norma ISO/IEC 14443-3: pętla antykolizyjna i identyfikacja kart 115
Kurs RFID z PN532, rozdział 3.
MIFARE Classic: mapa pamięci, autentykacja, zapis i odczyt 122
Kurs RFID z PN532, rozdział 4.
MIFARE Classic: Crypto1, ataki i ochrona własnych kart 129

FELIETON

Jak (dobrze) projektować elektronikę użytkową? 136

Prenumerata 2
Od wydawcy 3

Your
B2B
partner

Tak! Uproszczenie zamówień technicznych. Z Conrad.

Dopasowane rozwiązania e-Procurement



conrad.pl/tak-z-conrad

All parts of success

CONRAD

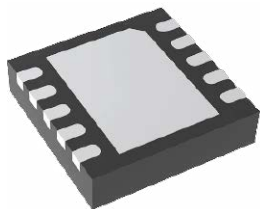
eprasa.pl 2f27e6d5ba

Nowe podzespoły

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

Bezpiecznik elektroniczny NIS5132

NIS5132 stanowi bezpiecznik elektroniczny (eFuse) przeznaczony do użycia w rozwiązaniach, które zasilane są napięciem nominalnym 12 V. Jest to element pracujący przy napięciu wejściowym: 9...18 V, co pozwala na jego stosowanie przede wszystkim w obwodach zasilania płyt głównych komputerów. Podzespół wyposażono także w tranzystor NMOS charakteryzowany przez rezystancję w stanie włączenia $R_{DS(on)}$ wynoszącą typowo 44 mΩ. Tranzystor sterowany jest za pośrednictwem wbudowanej pompy ładunkowej, a dodatkowo uwzględniony jest ogranicznik prądu razem z blokadą podnapięciową (UVLO) aktywującą się przy zbyt niskiej wartości napięcia wejściowego.



Podzespół skutecznie ogranicza przepięcia, utrzymując napięcie wyjściowe w dopuszczalnym zakresie pracy, bez konieczności odłączania obciążenia. Komponent oferuje również zabezpieczenie termiczne działające w trybie zatraskowym lub z automatycznym załączaniem w przypadku, gdy temperatura elementu przekracza 175°C.

Układ umożliwia kontrolę szybkości narastania napięcia, co ogranicza prąd udarowy w momencie załączania obciążenia. Zabezpieczenie znajduje zastosowania szczególnie w systemach zarządzania zasilaniem oraz w urządzeniach pamięci masowej, względem której wymagana jest pełna ochrona przed zwarciami i przepięciami, przy równoczesnym podtrzymaniu ciągłości pracy.

www.onsemi.com



Kondensatory mocy z serii MKK AC Modular

Seria kondensatorów mocy MKK AC Modular została stworzona z myślą o pracy w szerokiej gamie obwodów prądu przemiennego (AC). W podzespołach serii

stosuje się dielektryk z metalizowanej folii polipropylenowej, który charakteryzuje się zdolnością do regeneracji po dokonaniu przebicia. Dzięki temu, elementy wyróżniają się niezawodnością w pełnym zakresie aplikacji, takich jak: kompensacja mocy biernej lub filtracja harmonicznych w napięciu zasilania.

W zależności od modelu kondensatory MKK AC Modular pracują przy napięciu znamionowym w zakresie: 230...1000 V AC. Są to elementy oferowane w szerokim przedziale pojemności, które mogą być łączone w zestawy o większej mocy, przy równoczesnym zachowaniu równomiernego rozkładu obciążenia. Nie zapomniano również o uwzględnieniu trwałych obudów. Natomiast zastosowany dielektryk minimalizuje straty mocy i nagrzewanie, co warunkuje stabilność parametrów roboczych nawet w wymagających warunkach przemysłowych.

Wszystkie kondensatory serii przystosowane są do pracy w temperaturze otoczenia nieprzekraczającej 55°C. Ich prąd znamionowy zależy od warunków eksploatacji, a zastosowania obejmują przede wszystkim przekształtniki energoelektroniczne. Elementy sprawdzają się również w filtrach pasywnych, w przypadku których wymagane jest zapewnienie niezakłóconej pracy przy obciążeniach nieliniowych.

www.tdk-electronics.tdk.com

Czujnik drgań PKGM-210D-R

W ofercie Murata Manufacturing dostępny jest czujnik PKGM-210D-R umożliwiający detekcję drgań typowych dla uszkodzeń podzespołów wirujących.



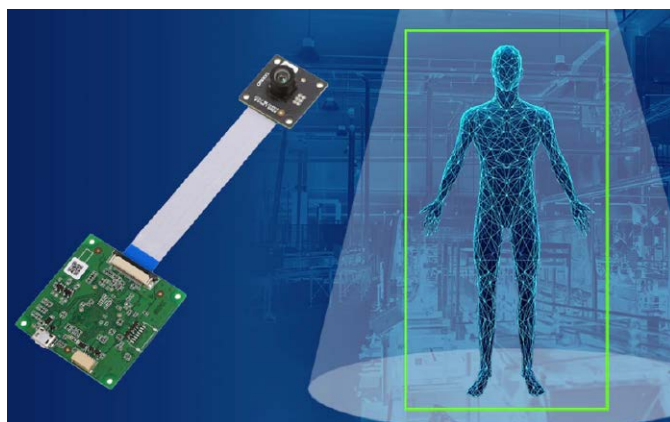
Dotyczy to przede wszystkim silników oraz łożysk, które szczególnie w przypadku niedostatecznego smarowania są źródłem drgań o częstotliwościach nieprzekraczających 20 kHz. Ich wykrywanie bywa utrudnione ze względu na szereg czynników, dlatego dotychczas opierało się ono głównie na metodach akustycznych. Zastosowanie czujnika PKGM-210D-R znacząco upraszcza ten proces, tym bardziej, że został uwzględniony przetwornik piezoelektryczny o paśmie roboczym: od 5 do 20 kHz. W rezultacie podzespół znajduje zastosowanie w diagnostyce predykcyjnej sprzętu, umożliwiając wykrywanie nawet bardzo subtelnych wibracji.

Prócz przetwornika czujnik PKGM-210D-R zawiera także dedykowany sterownik, którego obsługa nie sprawia trudności. Zintegrowany sensor temperatury umożliwia pełną kompensację pomiarów oraz bieżącą diagnostykę

komponentu. Natomiast wyniki pomiarów to wartości napięcia wyprowadzanego na wyjście podzespołu. Zakres pomiarowy PKGM-210D-R wynosi $\pm 10,2$ G, a typowa czułość przetwarzania osiąga 44 mV/G. Element zasilany jest napięciem stałym: 3...5,2 V. Zakres temperatury pracy mieści się w przedziale: od -20 do 85°C i dodatkowo rozwiązanie charakteryzuje średni pobór prądu równy 3 mA.

Obudowa czujnika o wymiarach: $5 \times 5 \times 3,5$ mm pozwala na montaż powierzchniowy elementu oraz uproszczoną jego integrację w szerokiej gamie aplikacji. Zastosowanie podzespołu sprzyja wczesnemu wykrywaniu usterek, ogranicza ryzyko nieplanowanych przestojów oraz umożliwia optymalizację harmonogramów serwisowych. Co więcej, przyczynia się również do przedłużenia żywotności elementów wirujących oraz równocześnie eliminuje konieczność wprowadzania zmian w urządzeniach opartych na czujniku PKGM-210D-R, co pozwala na szybkie jego wdrożenie przy obniżonym nakładzie pracy, bez ponoszenia dodatkowych kosztów z tym związanych.

www.murata.com



Moduł kamery B5T-007003 do detekcji obecności człowieka

Moduł B5T-007003 został wyposażony w miniaturową kamerę umożliwiającą wykrywanie obecności człowieka niezależnie od jego pozycji, zarówno stojącej, pochyłonej, kucznej, jak i podczas ruchu. Rozwiązanie składa się z dedykowanej jednostki obliczeniowej oraz toru akwizycji obrazu, a jego obsługa sprowadza się do przesyłania prostych poleceń z poziomu urządzenia nadrzędnego.

Rzeczywisty zakres detekcji zależy od sposobu montażu. Instalacja modułu na wysokości: 3...3,8 m, w orientacji pionowej zapewnia pole widzenia skorelowane z odległością od obserwowanej osoby. Z kolei montaż pod kątem 45° umożliwia lepsze dopasowanie pola detekcji do geometrii stanowiska pracy. Zastosowany algorytm pozwala na wykrywanie obecności człowieka, przy jednoczesnym ignorowaniu obiektów technicznych, co obniża ryzyko błędnych wskazań wynikających z obecności maszyn, źródeł ciepła, czy odbić światła.

Moduł działa poprawnie niezależnie od rodzaju noszonej odzieży roboczej, w tym odzieży ochronnej. Parametry pracy podzespołu konfiguruje się przy pomocy dedykowanego narzędzia instalacyjnego, które pozwala na m.in. wybór orientacji kamery, wysokości montażu, kierunku detekcji oraz zdefiniowanie stref wyłączonych poprzez maskowanie konkretnych obszarów obrazu.

Rozwiązanie przeznaczone jest do zastosowań w aplikacjach przemysłowych jako element układów sterowania służących do nadzoru obecności operatora. Moduł nie spełnia wymagań norm bezpieczeństwa funkcjonalnego i nie może być traktowany jako jedyny środek ochrony. W związku z tym konieczne jest stosowanie dodatkowych zabezpieczeń zgodnych z obowiązującymi przepisami i przeprowadzanie weryfikacji działania w rzeczywistych warunkach eksploatacyjnych.

<https://components.omron.com>



Układ scalony XDM700-1 do precyzyjnych pomiarów prądu i napięcia

W układzie XDM700-1 pomiary prądu i napięcia prowadzone są z wykorzystaniem przetworników analogowo-cyfrowych (ADC) o rozdzielczości 12 bitów, zapewniających niepewność pomiarową na poziomie: $\pm 0,4\%$ – dla napięcia, oraz $\pm 0,75\%$ – dla prądu. Wyniki pomiarów uzyskiwane przez układ udostępniane są zarówno poprzez analogowe wyjście IMON, jak i cyfrowy interfejs PMBus 1.3,

REKLAMA

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

w ramach którego transmisje danych realizowane są z częstotliwością do 1 MHz. Element dostarcza również dane dodatkowe, obejmujące wartości minimalne lub maksymalne, oraz pozwala na definiowanie progów ostrzegawczych dla wielkości przez niego monitorowanych w czasie rzeczywistym.

Oprócz podstawowych funkcji pomiarowych podzespół oferuje obliczanie i raportowanie mocy. W połączeniu z компактowymi wymiarami sprawia to, że XDM700-1 sprawdza się wszędzie tam, gdzie wymagane jest wierne i równocześnie nieprzerwane monitorowanie parametrów zasilania. Element dostarczany jest w obudowie VQFN-24 o wymiarach: 4×4 mm i pracuje w zakresie temperatury: od -40 do 125°C. Znajduje on zastosowanie przede wszystkim w serwerach lub urządzeniach telekomunikacyjnych, ale również wykorzystywany jest w systemach robotycznych, zasilaczach oraz instalacjach opartych na odnawialnych źródłach energii (OZE). Układ XDM700-1 należy do rodziny XDP dedykowanej obwodom zasilania w aplikacjach o wysokiej gęstości mocy. Rodzina ta obejmuje pełne spektrum zastosowań: od elektroenergetyki, aż po centra danych, gdzie decydującą rolę odgrywa wysoka sprawność pracy wraz z jej niezawodnością.

www.infineon.com

Antena odbiorcza DLFPV-M3-R

DLFPV-M3-R to antena odbiorcza, przeznaczona do pracy w paśmie częstotliwości: 5,725...5,85 GHz, która nie wymaga stosowania układów dopasowania na wejściu. Zysk energetyczny podzespołu nie przekracza 2 dBi, co stanowi kompromis pomiędzy kierunkowością, a równomiernością jego charakterystyki promieniowania, która w praktyce jest dookólna i pozwala na stabilny odbiór sygnału przy szerokim kącie pokrycia. Ma to decydujące znaczenie w systemach, w których orientacja nadajnika względem odbiornika ulega częstym zmianom. Element wyposażony jest w złącze SMA i dodatkowo cechuje go impedancja wejściowa 50 Ω. Współczynnik fali stojącej (WFS) nie przekracza wartości 2, przy wysokości anteny 8,8 cm oraz jej średnicy 3,5 cm. Kompaktowe wymiary elementu zdecydowanie ułatwiają integrację DLFPV-M3-R z przeróżnymi rozwiązaniami. Antena przeznaczona jest dla zastosowań wymagających utrzymania łączności radiowych o niedużych opóźnieniach transmisji, w tym dla systemów sterowania oraz obserwacji, w których przesyłane obrazy spełniają funkcję informacji zwrotnej tak dla operatorów, jak i urządzeń nadrzędnych, zapewniając równocześnie wysoką niezawodność oraz odporność na zmienne warunki pracy.

www.mikroe.com



Miniaturowy zegar atomowy ICPT-1

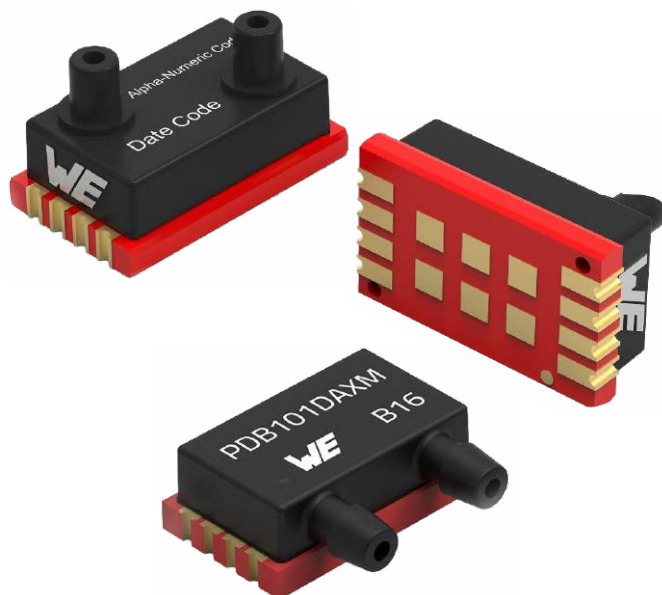
ICPT-1 to miniaturowy zegar atomowy, który zaprojektowano z myślą o wysokiej stabilności i niezawodności działania. Zastosowany w nim laser, inicjujący przejścia kwantowe, skutecznie eliminuje ograniczenia typowe dla klasycznych rozwiązań. W rezultacie podzespół charakteryzuje się niższym poborem mocy oraz zwiększoną stabilnością parametrów w długim okresie pracy.

Element pozwala na synchronizację czasu poprzez wyjście 1 PPS, zachowując ciągłość działania dzięki trybowi autonomicznemu (holdover), który ogranicza dryft częstotliwości. Precyzyjna regulacja przytoczonej wielkości, a także konfiguracja oraz odczyt parametrów zegara realizowane są przy pomocy interfejsu UART. Dodatkowo podzespół wyposażony został w licznik czasu rzeczywistego (Time-of-Day counter) o rozdzielczości 1 sekundy.

Do kluczowych parametrów ICPT-1 należą: napięcie zasilania 3,3 V i średni pobór prądu poniżej 0,5 A. Element oferuje także nominalną częstotliwość wyjściową 10 MHz przy tolerancji $\pm 0,05$ ppb. Stabilność krótkookresowa wynosi 0,09 ppb (przy $\tau=1$ s). Natomiast tempo starzenia nie przekracza 0,03 ppb w ciągu doby, w temperaturze 25°C. Zegar zamknięty jest w kompaktowej obudowie o wymiarach 45×36×14,5 mm, przystosowanej do pracy w zakresie temperatury: od -40 do 70°C.

Dzięki powyższym własnościom, komponent może spełniać rolę precyzyjnego źródła odniesienia w systemach wymagających wysokiej spójności czasowo-częstotliwościowej. Zegar atomowy ICPT-1 znajduje zastosowanie przede wszystkim w łączności satelitarnej, nawigacji oraz aplikacjach wymagających ścisłej synchronizacji – ze szczególnym uwzględnieniem systemów finansowych, energetycznych i telekomunikacyjnych.

www.iqdfrequencyproducts.com



Czujnik różnicy ciśnień WSEN-PDMS

Czujnik WSEN-PDMS przeznaczony jest do pomiaru różnic ciśnienia m.in. w układach automatyki i systemach wbudowanych. Niepewność pomiarowa podzespołu wynosi ± 1 mbar i równocześnie to element, który wyróżnia rozdzielczość 16 bitów. Jest on zasilany napięciem z zakresu: 3...5,5 V umożliwiającym bezpośrednią współpracę z szeroką gamą rozwiązań opartych na mikroprocesorach.

Na potrzeby komunikacji z otoczeniem w układzie przewidziane są interfejsy: I²C i SPI. Zwracane przez czujnik wyniki pomiarowe stanowią wartości skompensowane, w których uwzględnieniu podlega wpływ temperatury na podzespół. Parametry pracy komponentu można obserwować na bieżąco, a dane z niego wysyłane uzupełnia się o sumę kontrolną CRC zwiększającą niezawodność transmisji w ten sposób, że możliwe staje się wykrycie kilku błędów o charakterze losowym.

Zakres temperatury roboczej WSEN-PDMS mieści się w zakresie: od -25 do 85°C , co pozwala na zastosowanie podzespołu w standardowych środowiskach przemysłowych. Element bazuje na wypróbowanych rozwiązaniach, uwzględnionych wcześniej w czujniku WSEN-PDUS. Dzięki temu, komponent z powodzeniem sprawdza się np. w systemach HVAC, detektorach nieszczelności instalacji, automatyce przemysłowej i aparaturze medycznej, w tym w inhalatorach. Łącznie oferowane są dwa warianty WSEN-PDMS. Jeden z nich obejmuje króćce pionowe, dopasowane do adapterów ciśnienia oraz kolektorów pomiarowych. Natomiast drugi wariant wyposażony jest w króćce poziome, które są ząbkowane i umożliwiają bezpośrednie podłączanie się do przewodów elastycznych.

www.we-online.com

REKLAMA



arm KEIL MDK v6

Nowa generacja narzędzi dla systemów embedded

- **VS Code + CMSIS - Toolbox**
na Windows, Linux i macOS
- **Arm Compiler 6 (LLVM/Clang)**
i debugowanie Cortex-M
- **Integracja z CI/CD**
i narzędziami automatyzacji

Kup MDK v6 - 5% rabatu na wersję Professional

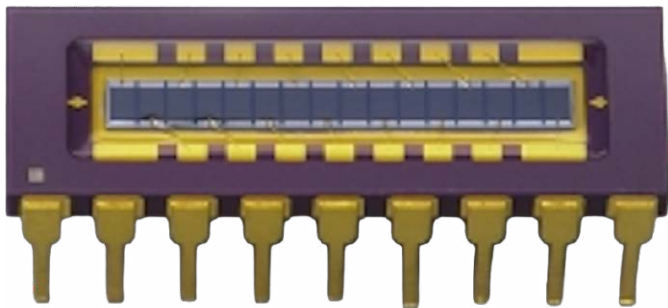
Dowiedz się więcej na www.ccontrols.pl



Computer Controls Sp. z o.o.

Bielsko-Biała, ul. Bystrzańska 94

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



Matryca fotodiod krzemowych S4111-16Q

S4111-16Q to matryca krzemowych fotodiod zamknięta w ceramicznej obudowie DIP o 18 wyprowadzeniach. Element składa się z 16 elementów detekcyjnych o wymiarach: 0,9×1,45 mm każdy, rozlokowanych linowo w jednym rzędzie. Taka konfiguracja umożliwia prowadzenie wielokanałowych pomiarów promieniowania: od nadfioletu, po bliską podczerwień (190...1100 nm). Maksymalna czułość matrycy przypada w okolicach długości fali 960 nm. Konstrukcja układu została zoptymalizowana pod kątem uzyskania wysokiego stosunku sygnału do szumu, co ma szczególne znaczenie w precyzyjnych aplikacjach pomiarowych.

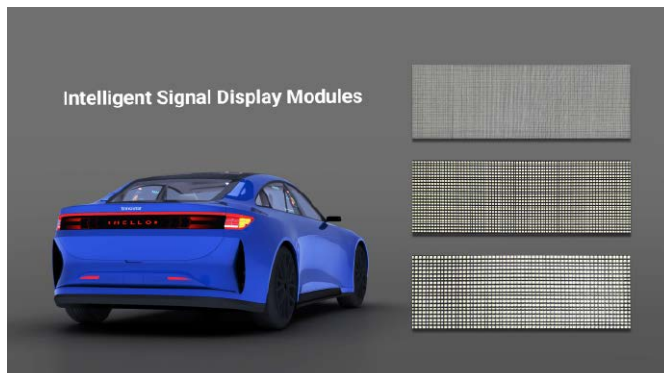
Maksymalne napięcie polaryzacji wstecznej wynosi -15 V. Prąd ciemny pojedynczej fotodiody nie przekracza 5 pA, co pozwala ograniczać wpływ szumów na tor pomiarowy. Pojemność złącza fotodiody na poziomie 200 pF wpływa na właściwości dynamiczne układu, w tym jego stałą czasową. Czas narastania sygnału nie przekracza 0,5 μs, co przekłada się na szybką odpowiedź matrycy na zmiany natężenia promieniowania. Wszystkie podane parametry odnoszą się do temperatury 25°C oraz pojedynczego elementu detekcyjnego.

Konstrukcja matrycy minimalizuje przesłuchy między sąsiednimi fotodiodami, poprawiając separację sygnałów. Ponadto duża powierzchnia aktywna podwyższa efektywność detekcji, a możliwość użycia filtrów optycznych pozwala dostosować charakterystykę układu do specyficznych wymagań aplikacyjnych. Dodatkową zaletą jest brak konieczności stosowania chłodzenia, co upraszcza integrację i zmniejsza wymagania eksploatacyjne, przy jednoczesnym zachowaniu wysokiej czułości i stabilnej pracy w szerokim zakresie warunków środowiskowych.

www.hamamatsu.com

Moduły microLED do zastosowań motoryzacyjnych

Technologia microLED należy do najważniejszych kierunków rozwoju współczesnej optoelektroniki. Choć początkowo była ona wykorzystywana tylko w systemach wyświetlania, obecnie znajduje zastosowanie



również w elektronice samochodowej oraz optycznych systemach transmisji danych o wysokiej przepustowości.

Przykładem postępu w tej dziedzinie są moduły microLED firmy Ennostar, które zapewniają wyższą sprawność emisji światła niż rozwiązania konkurencyjne. W przypadku barwy niebieskiej i zielonej wzrost ten wynosi 10...15%, natomiast dla światła czerwonego jest to nawet 90%. Ma to miejsce przy konsekwentnym udoskonalaniu półprzewodnikowych, co przekłada się na wysoką stabilność parametrów oraz lepsze dostosowanie modułów do produkcji wielkoseryjnej.

Kluczowym elementem rozwoju modułów microLED jest także miniaturyzacja – zmniejszenie wymiarów diod w nich zawartych o 40%. Pozwala to nie tylko obniżyć koszty produkcji, lecz również zwiększyć rozdzielczość wyświetlaczy opartych na modułach oraz poprawić ich efektywność świetlną. Tak rozumiane moduły znajdują zastosowanie m.in. w wyświetlaczach pokładowych, systemach projekcyjnych typu HUD, a nawet reflektorach adaptacyjnych, gdzie kluczowe znaczenie mają wysoka jasność, kontrast i energooszczędność.

Równolegle technologia microLED rozwijana jest z myślą o systemach transmisji optycznej wykorzystywanych w zaawansowanych aplikacjach. Moduły firmy Ennostar mogą pełnić funkcję źródeł światła w wielokanałowych układach transmisyjnych, umożliwiając szybki transfer danych na krótkich dystansach – rzędu kilku metrów. Ich wysoka sprawność oraz możliwość pracy w podwyższonych temperaturach czynią je szczególnie przydatnymi w wymagających miejscach pracy. Dalszy ich rozwój będzie zależny przede wszystkim od optymalizacji procesów produkcji, zwiększenia jednorodności emisji oraz redukcji kosztów wytwarzania. Istotne znaczenie będzie miało także opracowanie efektywnych metod integracji z układami sterującymi, co przełoży się na pełną niezawodność modułów oraz szersze ich zastosowania – zwłaszcza w systemach transmisji optycznej Li-Fi w celu szybkiej oraz odpornej na zakłócenia komunikacji w środowiskach o dużym zagęszczeniu urządzeń i ograniczonej dostępności pasma radiowego.

www.ennostar.com



Akcelerometry RECOVIB

Seria RECOVIB uwzględnia czujniki MEMS stworzone z myślą o pomiarze przyspieszeń liniowych w szerokim spektrum aplikacji. Dzięki specjalnym układom kondycjonowania sygnału, w sposób pełny ogranicza się wpływ zakłóceń, co przekłada się na niezawodną i stabilną transmisję danych, nawet na większe odległości. Dostępność różnych wariantów wyjść, w tym szczególnie odpornego na zakłócenia wyjścia prądowego, upraszcza integrację czujników z różnymi systemami. Podzespoły sprawdzają się doskonale zwłaszcza w trudniejszych środowiskach pracy, dla których znaczenie ma odporność na cięższe warunki eksploatacji.

Konstrukcja akcelerometrów RECOVIB obejmuje galwaniczną separację toru sygnałowego, która gwarantuje to, że nie są wymagane dodatkowe elementy izolacyjne. W zależności od wariantu wykorzystuje się pojemnościową lub piezorezystancyjną metodą pomiarów, a oferta czujników obejmuje wiele klas podzespołów, które różnią się poziomem szumu i rozdzielczością. Warianty standardowe wyróżniają się gęstością szumu poniżej $50 \mu\text{g}/\sqrt{\text{Hz}}$. Z kolei modele o podwyższonej rozdzielczości osiągają wartości poniżej $10 \mu\text{g}/\sqrt{\text{Hz}}$, a warianty specjalne schodzą z wymienionym parametrem do wartości mniej niż $2 \mu\text{g}/\sqrt{\text{Hz}}$. Dostępne są także czujniki o paśmie przenoszenia 24 kHz, przewidziane do kompleksowej diagnostyki maszyn.

Obudowa przyrządów przystosowana jest do wymagających warunków środowiskowych. Stopień ochrony IP67 zapewnia odporność na pyły i niewrażliwość na krótkotrwałe zanurzenie w wodzie. Dodatkowo wykorzystaniu podlegają obudowy aluminiowe lub stalowe, w tym obudowy wykonane ze stali nierdzewnej 316L dedykowanej do zastosowań o podwyższonej korozyjności.

Akcelerometry RECOVIB dostarczane są w konfiguracjach: jedno-, dwu- oraz trójosiowych, co pozwala na ich precyzyjne dopasowanie do specyfiki aplikacji. Czujniki znajdują zastosowanie m.in. przy monitorowaniu drgań maszyn lub w diagnostyce układów i kontroli procesów technologicznych oraz są rekomendowane do użycia w celu nadzoru konstrukcji, a także w systemach tłumienia drgań i aplikacjach związanych z energetyką wiatrową.

<https://micromega-dynamics.com>



Czujniki węglowodorów gazowych C1-C6

Opracowane czujniki wykorzystują metodę niedyspersyjnej absorpcji promieniowania podczerwonego, w której stężenie węglowodorów gazowych C1-C6 wyznaczane jest na podstawie stopnia pochłaniania promieniowania o charakterystycznej długości fali. Układy optyczne podzespołów uwzględniają źródło promieniowania IR razem z precyzyjnym detektorem, który pozwala na selektywną identyfikację analizowanych gazów. Wszystkie pomiary przeprowadza się bezkontaktowo, a dodatkowo czujniki wyróżnia trwałość, przewyższająca standardy spotykane w konkurencyjnych rozwiązaniach tej klasy.

Aby zapewnić stabilność wskazań, czujniki firmy Honeywell wyposażono w mechanizmy kompensujące wpływ temperatury oraz starzenia się komponentów. W warunkach stabilnego przepływu gazów uzyskiwane wyniki są powtarzalne i spełniają wymagania szerokiej gamy aplikacji przemysłowych. Natomiast zakres pomiarowy oraz rozdzielczość zależą od konkretnych wariantów czujników. Integracja podzespołów z systemami automatyki jest uproszczona i staje się możliwa, dzięki standardowym interfejsom komunikacji.

Zastosowana metoda pomiarów eliminuje konieczność przeprowadzania kalibracji z użyciem gazów odniesienia (wzorcowych). Czujniki przeznaczone są do użycia w środowiskach, w których istnieje ryzyko gromadzenia się niebezpiecznych gazów. Sensory współpracują m.in. z systemami alarmowymi, wentylacyjnymi oraz systemami odcinania zasilania. W przypadku przekroczenia określonego progu stężenia generowany jest sygnał wyjściowy, który może sterować elementami wykonawczymi. Oprócz tego elementy cechują się obniżonym poborem mocy, a po uruchomieniu przeprowadzana jest procedura autokalibracji, po której czujniki przechodzą do pracy ciągłej. Są to podzespoły odporne na zakłócenia elektromagnetyczne (EMI) – zgodnie z obowiązującymi standardami, co bezpośrednio wpływa na bezpieczeństwo i działanie systemów detekcji gazów oraz związanych z nimi układów, w tym zintegrowanych zabezpieczeń.

<https://automation.honeywell.com>



Bateria litowa wielokrotnego ładowania ML414H

ML414H to niewielka bateria przeznaczona do aplikacji o obniżonym poborze energii. Jej napięcie znamionowe wynosi 3 V, a zalecany zakres napięcia ładowania zawiera się między: 2,8 a 3,1 V. Pojemność ogniwa to 1 mAh, przy rezystancji wewnętrznej 600 Ω oraz maksymalnym prądzie rozładowania do 5 μ A. Dzięki tym parametrom, bateria sprawdza się doskonale w bogatej gamie zastosowań, w tym w rozwiązaniach profesjonalnych, wliczając w to układy podtrzymania pamięci oraz moduły zegara czasu rzeczywistego (RTC).

Kompaktowe wymiary ML414H, określone przez średnicę 4,8 mm i wysokość 1,4 mm, oraz ciężar około 70 mg sprawiają, że bateria nadaje się wprost do urządzeń o ograniczonej przestrzeni montażowej. Dodatkowym atutem ogniwa jest jego szczelna obudowa, która w standardowych warunkach pracy skutecznie wyciekiem elektrolitu, z łatwością zwiększając bezpieczeństwo użytkowania podzespołu.

Aby zapewnić prawidłowe i bezpieczne działanie baterii, należy przestrzegać podstawowych zasad eksploatacji. Przekroczenie dopuszczalnych parametrów – takich jak: napięcie ładowania, prąd oraz temperatura pracy, może doprowadzić do przegrzewania, uszkodzenia baterii lub skrócenia jej żywotności. Baterię należy przechowywać również w kontrolowanych warunkach, z dala od wilgoci i bezpośredniego nasłonecznienia. Wskazane jest także unikanie kontaktu z elementami przewodzącymi oraz przechowywanie ML414H w oryginalnym opakowaniu, co redukuje ryzyko przypadkowego zwarcia. Podczas użytkowania szczególnie ważne jest zachowanie prawidłowej polaryzacji, unikanie przeciążeń i ochrona baterii przed uszkodzeniami. Ich wystąpienie może stanowić poważne zagrożenie dla zdrowia użytkownika, prowadząc do poważnych konsekwencji zwłaszcza na skutek eksplozji lub wycieku elektrolitu grozących poważnymi obrażeniami.

www.sii.co.jp



Aktywny przedłużacz optyczny ze złączem USB 3.2 Gen 2

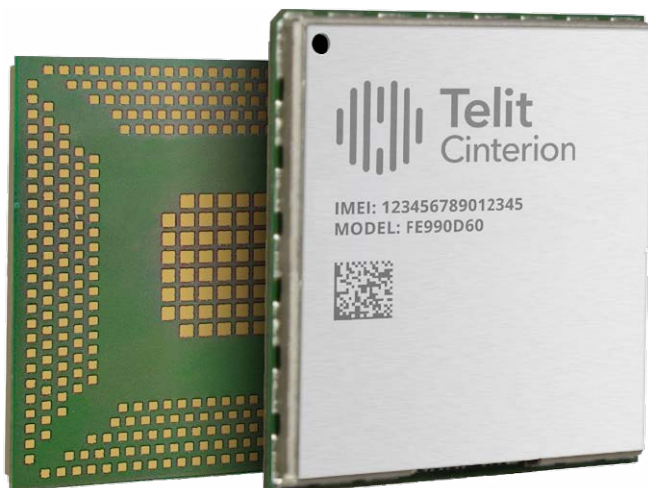
Przedłużacz firmy Amphenol pozwala równocześnie na transmisję danych i dostarczanie zasilania do urządzeń, stanowiąc praktyczne rozwiązanie w sytuacjach, gdy tradycyjne przewody miedziane nie spełniają wymagań ze względu na ograniczoną długość. Dzięki zastosowaniu toru optycznego, eliminowane są problemy wynikające z tłumienia sygnałów, gwarantując stabilną komunikację na odległość do 15 m – bez pogorszenia jej kluczowych parametrów.

Rozwiązanie jest zgodne ze standardem USB 3.2 Gen 2 i zapewnia przepustowość transmisji danych do 10 Gb/s. Przedłużacz umożliwia w szczególności przesyłanie obrazu o bardzo wysokiej rozdzielczości: nawet do 8K przy częstotliwości odświeżania 60 Hz. Dodatkowo obsługa technologii Power Delivery (PD) oznacza możliwość, by podłączone urządzenia zasilają mocą do 60 W – przy zachowaniu kompatybilności wstecznej ze standardem USB 2.0.

Konstrukcja mechaniczna przedłużacza sprawdza się idealnie w wymagających środowiskach. Przede wszystkim dostępne jest złącze gwintowane, zgodne z normą MIL-DTL-38999, które zapewnia pewne i stabilne połączenia, nawet w miejscach narażonych na silne drgania i obciążenia mechaniczne. Przedłużacz spełnia również wymogi, które dotyczą: stopnia ochrony IP68 oraz klasy wytrzymałości IK06. Dodatkowo produkt cechuje niewrażliwość na zakłócenia elektromagnetyczne (EMI) – w tym względzie spełnia on wytyczne normy MIL-STD-810H.

Dzięki powyższym właściwościom, przedłużacz znajduje swoje zastosowania w szerokim spektrum aplikacji: od przemysłu, aż po sektor militarny. Rozwiązanie sprawdza się doskonale wszędzie tam, gdzie wymaga się ciągłej wymiany danych i niezawodnego zasilania, zwłaszcza w trudno dostępnych lokalizacjach. W przemyśle produkt stosowany jest m.in. na liniach produkcyjnych, placach budowy lub w warsztatach. Natomiast w zastosowaniach wojskowych przedłużacz staje się nieodzowny podczas działań „w terenie” – jego trwałość i niezawodność mają decydujące znaczenie, niezależnie od czynników takich jak: wilgoć i pyły i wysokie temperatury, które stanowią poważne zagrożenie dla życia i zdrowia człowieka.

www.amphenol.com



Moduł komunikacji FE990D60

Moduł FE990D60 to zaawansowany komponent w obudowie LGA, opracowany z myślą o transmisji danych w sieciach komórkowych 5G. Dzięki wstecznej kompatybilności z technologiami LTE oraz WCDMA, podzespół zapewnia ciągłość połączenia także w obszarach pozbawionych zasięgu 5G. Urządzenie wyposażono w cztery wejścia antenowe 4G/5G oraz w wejście antenowe dla systemów nawigacji satelitarnej, które obsługują pasma częstotliwości L1 i L5. Umożliwia to równoczesną realizację funkcji komunikacyjnych i precyzyjnej lokalizacji.

Zintegrowany procesor ARM Cortex-A55, taktowany zegarem 2,2 GHz, pozwala na efektywne przetwarzanie danych oraz uruchamianie aplikacji bez konieczności korzystania z zewnętrznych jednostek obliczeniowych. Moduł FE990D60 pracuje pod kontrolą systemu operacyjnego opartego na Linux (OpenWRT), który zapewnia stabilne oraz

elastyczne środowisko do tworzenia i uruchamiania aplikacji użytkownika. Wysoką przepustowość transmisji osiągnięto dzięki zastosowaniu agregacji nośnych oraz modulacji 256-QAM. Oferowany jest również szeroki zestaw interfejsów, w tym: USB 3.1 Gen 2, USB 2.0, PCIe, UART oraz USXGMII, co ułatwia znacząco integrację z różnorodnymi systemami i urządzeniami.

Moduł obsługuje także dwie karty SIM, a także umożliwia współpracę z maksymalnie trzema odbiornikami Wi-Fi lub dwoma urządzeniami Ethernet. Dodatkowo FE990D60 oferuje: dwa interfejsy USXGMII, trzy interfejsy UART oraz do trzech interfejsów PCIe. Kompaktowa obudowa LGA o wymiarach: 52×52×2,9 mm nadaje się do montażu powierzchniowego (SMT). Zakres temperatur pracy od -40 do 85°C pozwala na użycie układu w wymagających środowiskach przemysłowych. Element wspiera mechanizm FOTA (Firmware Over-The-Air), umożliwiając zdalne aktualizacje oprogramowania, oferując warstwę sterowników kompatybilną z systemami operacyjnymi: Windows 11 i Linux.

Dzięki integracji funkcji komunikacyjnych, obliczeniowych oraz interfejsowych w jednym układzie, FE990D60 znajduje zastosowanie np. w systemach bezprzewodowego dostępu do internetu oraz do transmisji wideo. Moduł stanowi zalecane rozwiązanie dla aplikacji IoT i telekomunikacyjnych, wliczając w to routery i bramy sieciowe. Zapewnia on wydajność oraz prostą integrację z różnego rodzaju systemami, w tym z systemami wbudowanymi, coraz chętniej stosowanymi, nawet tymi, które powstały dla najprostszych aplikacji.

www.telit.com

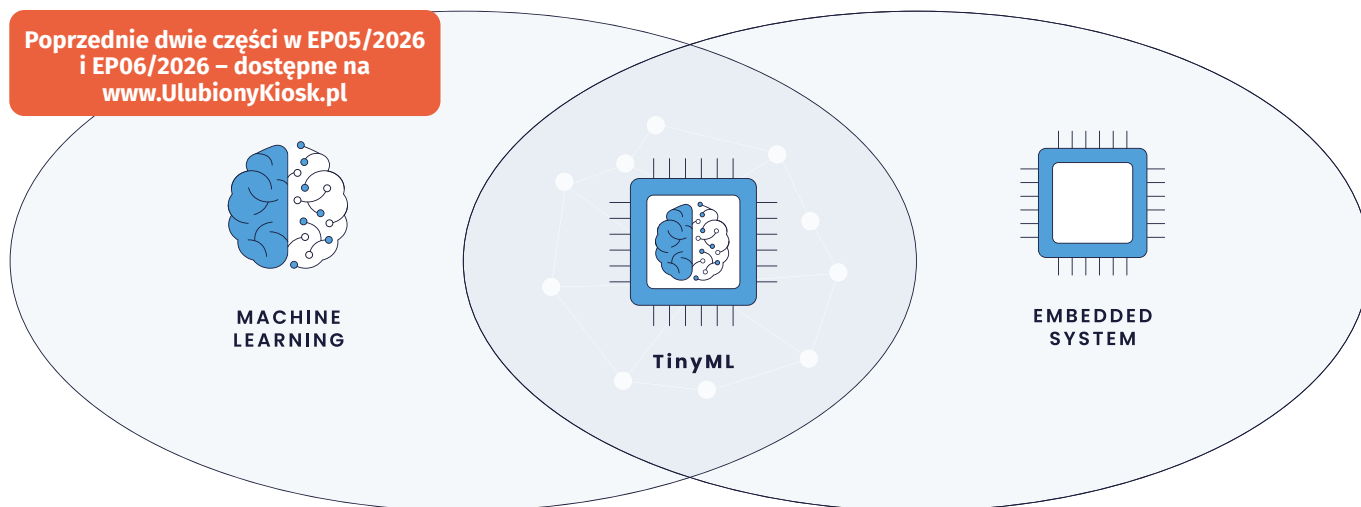
REKLAMA

świat radio
Magazyn wszystkich użytkowników eteru
KRÓTKOFALARSTWO CB RADIOTECHNIKA

przejrzyj i kupisz na stronie
www.ulubionykiosk.pl

Najnowsze doniesienia o aktywności Słonecznej
świat radio
Magazyn wszystkich użytkowników eteru
KRÓTKOFALARSTWO CB RADIOTECHNIKA
Summits On The Air
HCOM IC-7300MK2
Yaesu FTX-1 Field Optim
7-8/20

Poprzednie dwie części w EP05/2026
i EP06/2026 – dostępne na
www.UlubionyKiosk.pl



AI w pracowni elektronika konstruktora (3)

TinyML: duże uczenie w małym układzie

Deep Learning zrewolucjonizował sztuczną inteligencję – ale ma swoją cenę. Potrzebuje ogromnych zbiorów danych, specjalistycznego sprzętu i megawatów energii. TinyML to nie okrojony deep learning. To zupełnie inna filozofia: nie chodzi o to jak zrobić model jak najdokładniejszy, lecz jak zrobić model jak najbardziej użyteczny przy ekstremalnych ograniczeniach zasobów.

3.1 Czym TinyML jest, a czym nie jest

Definicja TinyML, chociaż pozornie prosta, kryje w sobie istotne niuanse. Formalna granica to urządzenia działające przy poborze mocy poniżej 1 mW, z pamięcią SRAM 32..512 kB i Flash poniżej 1 MB. Urządzenia te z reguły nie mają systemu operacyjnego, są jednotaktowe i unikają dynamicznej alokacji pamięci – bo po prostu nie mogą sobie na nią pozwolić. Cała inteligencja musi zmieścić się w firmware, działającym bezpośrednio na gołym metalu (bare-metal).

TinyML nie jest równoznaczny z prostą klasyfikacją ani z trywialnymi modelami. Badacze z MIT wykazali, że przy odpowiednim co-designie algorytmu i stosu systemowego możliwe jest uruchomienie modeli skali ImageNet na urządzeniach IoT – to, co jeszcze kilka lat temu wydawało się niemożliwe na sprzęcie z setkami kilobajtów pamięci. Hasło Lin, Zhu i Hana brzmi: dzisiejszy duży model jest jutrzejszym małym modelem. Granica TinyML przesuwana się wraz z rozwojem technik kompresji i sprzętowych akceleratorów.

TinyML to nowe pogranicze uczenia maszynowego. Wciskając modele deep learning w miliardy urządzeń IoT i mikrokontrolerów, rozszerzamy zakres zastosowań AI i umożliwiamy wszechobecną inteligencję.

– Lin J., Zhu L., Han S. et al.
Tiny Machine Learning: Progress and Futures,
IEEE CAS Magazine 2024

Kluczowe rozróżnienie, które inżynier musi mieć w głowie: TinyML to kategoria zastosowania, nie kategoria algorytmów. Ten sam algorytm Random Forest może być TinyML na STM32 i zwykłym Edge AI na Raspberry Pi. To sprzęt i budżet zasobów definiują przynależność do tej kategorii, nie architektura sieci.

3.2 Trzy fundamentalne wyzwania – i dlaczego występują jednocześnie

Każdy, kto pracował przy projekcie embedded AI, zna ten moment: model działa pięknie w środowisku testowym, a po wdrożeniu na urządzeniu zachowuje się nie tak jak

Tabela 6. Trzy jednocześnie wyzwania embedded AI. Opracowanie własne na podst.: Renesas/EE World 2025; Heydari i Mahmoud, Sensors 2025

Wyzwanie	Na czym polega	Dlaczego trudne w embedded
Wariacja danych	Dane rzeczywiste są szumione i zmienne – ten sam sygnał wygląda inaczej w różnych warunkach	Mały model uczy się na ograniczonej liczbie próbek, łatwo się przetrenuje do warunków laboratoryjnych
Detekcja w czasie rzeczywistym	System musi reagować natychmiast lub w ściśle określonym oknie czasowym	MCU ma ograniczoną moc obliczeniową; każda milisekunda opóźnienia ma znaczenie
Trójkąt ograniczeń: rozmiar–moc–cena	Forma, masa, energia i koszt jednocześnie ograniczają przestrzeń projektu	Zwiększenie dokładności modelu wymaga większego modelu, co zwiększa pamięć i pobór mocy

powinien. Za tym doświadczeniem stoją trzy wyzwania, które w świecie embedded AI występują zawsze jednocześnie i wzajemnie się wzmacniają.

Wariacja jest najtrudniej przewidywalnym z tych trzech problemów. Jeżeli chcemy wykrywać przysiady w urządzeniu noszonym, szybko okazuje się, że różni ludzie wykonują je w różny sposób – z różnymi prędkościami, amplitudami i wzorcami zmęczenia. Do tego dochodzi wariacja tła: czujnik wibracji zamontowany na urządzeniu przemysłowym rejestruje nie tylko działanie docelowej maszyny, ale także drgania przeniesione przez konstrukcję z sąsiadujących urządzeń. Model, który nie widział tej wariacji podczas treningu, nie będzie w stanie sobie z nią poradzić w terenie.

Trójkąt ograniczeń to matematyczny kompromis, z którego nie ma ucieczki. Zwiększenie częstotliwości próbkowania danych poprawia dokładność wykrycia, ale zwiększa wymagania pamięci podręcznej i obciążenie CPU. Zmniejszenie okna decyzyjnego przyspiesza odpowiedź, ale zmniejsza kontekst dla modelu. Zwiększenie precyzji wag poprawia jakość modelu, ale wymaga więcej pamięci i zwalnia operacje na MCU bez FPU. Każda decyzja projektowa przesuwają jeden z wierzchołków trójkąta kosztem pozostałych.

3.3 Pipeline TinyML od danych do firmware

W odróżnieniu od klasycznego ML, gdzie pipeline kończy się na serwerze, TinyML ma dodatkowe stadia specyficzne dla środowiska embedded. Generyczny przepływ pracy składa się z pięciu etapów, z których ostatnie dwa są unikalne dla tego paradygmatu.

Trening modelu odbywa się z reguły poza urządzeniem docelowym – na serwerze lub stacji roboczej z GPU. To normalne i pożądane: zasoby potrzebne do treningu są rzędy wielkości większe niż te dostępne na MCU. Wyjątek stanowi on-device training, omówiony szerzej poniżej. Punkt przekazania

między środowiskiem treningowym a docelowym następuje po etapie optymalizacji: model wychodzi ze środowiska ML (PyTorch, TensorFlow, scikit-learn) jako skompresowany, skwantyzowany artefakt gotowy do konwersji.

Etap konwersji jest często niedoceniany, a bywa najtrudniejszy. Konwertery takie jak TensorFlow Lite for Microcontrollers, ONNX Runtime czy dedykowane SDK producentów generują kod C/C++ lub binarny, który może być wkompiłowany w firmware. Na tym etapie ujawniają się problemy z precyzją numeryczną: operacje, które w środowisku treningowym wykonywane były na 32-bitowych liczbach zmiennoprzecinkowych, muszą być zreprezentowane przez 8-bitowe liczby całkowite. Sprawdzenie, że wynik klasyfikowania po konwersji nie różni się istotnie od oryginału, jest obowiązkowe przed zabezpieczeniem firmware.

3.4 Cztery klasy zastosowań TinyML – i co je łączy

Przegląd literatury pokazuje, że TinyML koncentruje się dziś wokół czterech głównych klas zastosowań. Nie są one

BENCHMARKI TINYML – MLPerf Tiny i MLMark

Jak mierzyć wydajność modeli TinyML w sposób porównywalny między platformami?

Branża odpowiedziała dwoma standardami benchmarkingowymi:

MLPerf Tiny (MLCommons/EEBMC) – cztery zadania referencyjne:

- **Keyword Spotting (KWS)** – rozpoznawanie słów kluczowych
- **Visual Wake Words (VWW)** – wykrywanie obecności osoby w kadrze
- **Image Classification** – klasyfikacja obrazów niskiej rozdzielczości
- **Anomaly Detection** – wykrywanie anomalii z danych sensorycznych

Modele referencyjne poniżej 250 kB. Metryki: latencja + energia na inferencję.

MLMark (EEBMC) – szersze pokrycie platform, standardowe punkty odniesienia dla wbudowanego ML; reprodukowalne wyniki bazowe.

Oba standardy są uzupełnione: MLPerf mierzy efektywność energetyczną, MLMark – dokładność i przyspieszenie względem linii bazowej.

1 Zbieranie danych sensory, etykietowanie, ground truth	2 Trening modelu serwer lub chmura, framework ML	3 Optymalizacja kwantyzacja, pruning, NAS	4 Konwersja TFLite/ONNX/vendor SDK	5 Rozlokowanie firmware, bare-metal, walidacja in-situ
---	--	---	--	--

3.1. Generyczny pipeline TinyML. Etapy 4–5 są specyficzne dla środowiska embedded. Opracowanie własne na podst.: Lin et al. 2024; Huang et al. 2025

Mapa zastosowań TinyML



3.2. Przekrój zastosowań TinyML we współczesnych urządzeniach elektronicznych

przypadkowe – wynikają bezpośrednio z charakterystyki środowisk, w których latencja, prywatność lub brak połączenia sieciowego wykluczają podejście chmurowe.

Urządzenia noszone (wearables) to historycznie pierwsza nisza TinyML. Smartwatche, opaski sportowe, plastry EKG, aparaty słuchowe – wszystkie muszą działać tygodniami lub miesiącami na małej baterii, reagować w czasie rzeczywistym i często przetwarzać dane biometryczne, które z powodów prywatności lub regulacyjnych nie mogą opuszczać urządzenia. Badania Heydari i Mahmoud pokazują, że latencja inferencji w zastosowaniach healthcare utrzymuje się ok. 100 ms dla typowych zadań diagnostycznych, z dokładnością w zakresie 80...99%. złożone diagnozy, jak wykrywanie zmian nowotworowych, wymagają 1–2 s nawet na sprofilowanym sprzęcie.

Keyword Spotting (KWS) to technicznie najtrafniejszy przykład dojrzałości TinyML. Rozpoznawanie słów budzących działa dziś na dedykowanych układach DSP o poborze mocy rzędu mikrowatów. Urządzenie nasłuchuje bez przerwy; model uruchamia pełne przetwarzanie dopiero po wykryciu słowa budzącego. To klasyczny wzorzec always-on z minimalnym zużyciem energii w stanie czuwania.

Detekcja anomalii w urządzeniach przemysłowych to obszar, gdzie TinyML wnosi konkretną wartość ekonomiczną. Czujnik wibracji zamontowany na silniku elektrycznym, sprężarce lub młynku może wykrywać znaki zbliżającego się uszkodzenia – łożysko z bieżnikiem do wymiany, niewyważone obciążenie – z wyprzedzeniem kilku dni lub tygodni. To utrzymanie predyktywne

Tabela 7. Klasy zastosowań TinyML z profilami wymagań. Opracowanie własne na podst.: Heydari i Mahmoud 2025; Lin et al. 2024; Tsoukas et al. 2024.

Klasa zastosowań	Przykładowe urządzenia	Kluczowy wymóg	Typowa latencja
Wearables/healthcare	smartwatch, opaska EKG, aparat słuchowy	Niski pobór mocy, prywatność	100 ms...2 s
Keyword Spotting (KWS)	głośniki smart, piloty, terminale IoT	Always-on, mikrowaty w czuwaniu	<10 ms
Detekcja anomalii	czujniki przemysłowe, utrzymanie predyktywne	Dokładność, ciągłe monitorowanie	0,18...300 ms
Rozpoznawanie gestów/HMI	automotive, terminale, AGD bez ekranu	Latencja <50 ms, odporność na szum	<50 ms
Wizja embedded	kamery bezpieczeństwa, kontrola jakości	Moc obliczeniowa, pamięć dla modelu	10 ms...1 s

(predictive maintenance) w praktyce: zamiast awarii, planowane postoje.

Rozpoznawanie gestów i interfejsy bez ekranu to czwarta duża klasa zastosowań, szczególnie istotna w kontekście HMI w automotive i przemyśle. TinyML na IMU pozwala urządzeniu interpretować ruchy dłoni, położenie ciała lub drganie jako komendy – bez kamery, bez mikrofonów, bez połączenia sieciowego. Badania wykazują, że akcelerometr taktowany z częstotliwością 4 kHz (versus standardowe 0,1 kHz w większości wearable) pozwala rozróżnić nie tylko gesty, ale także rodzaj trzymanego przedmiotu na podstawie drgań przenoszonych przez ciało.

3.5 On-device training: od statycznej inferencji do uczącego się urządzenia

Przez pierwsze lata TinyML oznaczał wyłącznie inferencję: model był trenowany gdzie indziej, wgrywany do firmware raz i działał w terenie bez możliwości samodzielnej adaptacji. To ograniczenie było akceptowalne, ale dalekie od ideału – urządzenie nie mogło dostosowywać się do zmieniających się warunków bez interwencji producenta.

On-device training zmienia to założenie. Zamiast statycznego modelu, urządzenie może aktualizować swoje wagi w odpowiedzi na nowe dane zebrane w terenie. Brzmi prosto, ale implementacja na MCU z 256 kB RAM jest inżynierskim wyzwaniem pierwszego rzędu – przede wszystkim dlatego, że klasyczne algorytmy uczenia (optymalizatory gradientowe: Adam, SGD) wymagają przechowywania gradientów dla każdego parametru sieci, co oznacza kilkukrotne powielenie rozmiaru modelu w pamięci roboczej.

Adhikary, Dutta i Dwivedi w swojej pracy TinyWolf z 2024 roku pokazali inne podejście: zamiast optymalizatorów gradientowych, wykorzystali metaheurystyczny algorytm Grey Wolf Optimizer, zmodyfikowany pod kątem zarządzania pamięcią. Kluczowa innowacja to podział ról populacji wilków (alfa, beta, delta, omega) na segmenty Flash i SRAM ze stronicowaniem – tylko aktywna część modelu jest w danym momencie w pamięci operacyjnej. Efekt: trening modelu głębokiej sieci neuronowej na mikrokontrolerze z zaledwie 256 kB RAM przy 71% oszczędności

On-Device Training – stan technologii w 2025 roku

CO JUŻ DZIAŁA:

- Trening na MCU z 256 kB RAM (TinyWolf, 2024)
- Transfer learning na urządzeniu – dostrajanie modelu do nowego kontekstu
- Online learning: aktualizacja wag w czasie pracy urządzenia
- Federated learning: wspólne uczenie wielu urządzeń bez centralizacji danych

AKTUALNE OGRANICZENIA:

- Trening pełny (od zera) na MCU: tylko proste modele, dużo czasu i energii
- Brakuje standardowych frameworków do on-device training na bare-metal
- Bezpieczeństwo aktualizacji modeli OTA nie jest ustandaryzowane
- Brak solidnych benchmarków dla systemów TinyML uczących się w terenie

KIERUNEK: przejście od statycznej inferencji do dynamicznego, adaptacyjnego uczenia.

(Źródła: Huang et al. Electronics 2025; Adhikary et al. IoT 2024)

pamięci względem klasycznych optymalizatorów. Zużycie energii podczas treningu dla najmniejszej konfiguracji mieści się w przedziale 0,035...0,157 mWh.

3.6 Frameworki i narzędzia: co jest dostępne dla inżyniera

Ekosystem narzędzi TinyML dojrzał w tempie, które jeszcze kilka lat temu byłoby trudne do przewidzenia. Dziś inżynier ma do dyspozycji kilka sprawdzonych ścieżek rozłokowania, a wybór między nimi zależy głównie od docelowej platformy sprzętu i wymagań wydajnościowych.

TensorFlow Lite for Microcontrollers to dziś standard de facto – z zastrzeżeniem, że jego siła jest również jego słabością. Szerokie wsparcie platform oznacza, że nie jest on zoptymalizowany pod żadną konkretną architekturę. Na procesorach ARM Cortex-M z Helium (Armv8.1-M) warto przejść do CMSIS-NN – biblioteki ARM zoptymalizowanej pod instrukcje SIMD tej architektury, która potrafi wyciągnąć znacznie więcej operacji na cykl zegara.

EdgeImpulse zajmuje szczególnie pozycję w tym ekosystemie jako platforma end-to-end: od zbierania danych przez przeglądarkę, przez trening modelu w chmurze, po automatyczne generowanie firmware dla wybranego urządzenia. W 2025 roku platforma rozszerzyła wsparcie o modele NVIDIA TAO dla urządzeń Jetson Orin, otwierając drogę

Tabela 8. Główne frameworki rozłokowania TinyML/Edge AI. Opracowanie własne na podst.: Tsoukas et al. 2024; Huang et al. 2025; Edge Impulse 2025

Framework/narzędzie	Producent	Docelowy sprzęt	Charakterystyka
TensorFlow Lite for MCU	Google/TFLite	ARM Cortex-M, ESP32, STM32	Najszerzej stosowany; modele .tflite; runtime w C++
ONNX Runtime (Mobile)	Microsoft	ARM, x86, RISC-V	Interoperacyjność między frameworkami ML
CMSIS-NN	ARM	Cortex-M (wszystkie)	Biblioteka niskopoziomowa; max. wydajność na ARM
Edge Impulse	Edge Impulse Inc.	Szerokie wsparcie MCU/MPU	Platforma end-to-end: dane, trening, deployment
NanoEdge AI Studio	ST Microelectronics	STM32 series	No-code dla detekcji anomalii; AutoML na MCU
DRP-AI TVM	Renesas (Edgecortex)	RZ/V series MPU	Optymalizacja pod akceleratori DRP-AI3; generuje kod C



do detekcji defektów wizualnych z algorytmem FOMO-AD nawet na skromniejszym sprzęcie embedded.

Podsumowanie

TinyML to nie okrojony deep learning, lecz autonomiczna filozofia projektowania systemów AI pod ekstremalne ograniczenia zasobów: pobór mocy <1 mW, Flash <1 MB, SRAM 32...512 kB. Trzy fundamentalne wyzwania – wariacja danych, detekcja w czasie rzeczywistym i trójkąt ograniczeń rozmiar–moc–cena – występują zawsze jednocześnie i wymagają kompromisów na każdym etapie projektu.

Generyczny pipeline TinyML obejmuje pięć etapów: zbieranie danych, trening (na zewnętrznym sprzęcie), optymalizacja (kwantyzacja + pruning),

konwersja do formatu firmware i rozlokowanie (deployment) z walidacją in-situ. Cztery dominujące klasy zastosowań to: wearables/healthcare (latencja ok. 100 ms, dokładność 80...99%), Keyword Spotting (always-on, mikrowaty), detekcja anomalii i rozpoznawanie gestów/HMI.

On-device training przełamuje statyczny model TinyML: urządzenia z 256 kB RAM potrafią już trenować głębokie sieci neuronowe w terenie, bez połączenia z chmurą. To zapowiedź kolejnego przełomu paradygmatu – od statycznej inferencji do adaptującego się urządzenia.

WM, JS

Źródła

- [1] Lin J., Zhu L., Chen W.-M., Wang W.-C., Han S. Tiny Machine Learning: Progress and Futures. IEEE Circuits and Systems Magazine 2024. arXiv 2403.19076.
- [2] Heydari S., Mahmoud Q.H. Tiny Machine Learning and On-Device Inference: A Survey. Sensors 2025, 25(10), 3191. doi: 10.3390/s25103191.
- [3] Adhikary S., Dutta S., Dwivedi A.D. TinyWolf – Efficient On-Device TinyML Training for IoT. Internet of Things 2024, 28, 101365. doi: 10.1016/j.iot.2024.101365.
- [4] Tsoukas V. et al. A Review on the Emerging Technology of TinyML. ACM Computing Surveys 2024, 56, 259. doi: 10.1145/3661820.
- [5] Somvanshi S. et al. From Tiny Machine Learning to Tiny Deep Learning: A Survey. ACM Comput. Surv. 2025, 58(7), 168.
- [6] Huang X., Wang H., Qin S., Tang S.-K. Embedded AI: A Comprehensive Literature Review. Electronics 2025, 14(17), 3468.
- [7] Edge Impulse. 2024 in Review – FOMO-AD, NVIDIA AI na urządzeniach edge, wsparcie Jetson Orin.

REKLAMA



Idealny wzmacniacz operacyjny, który oscyluje

Osiem powodów, dla których układ poprawny na schemacie ideowym traci stabilność na płycie – i co z każdym z nich zrobić

Wzmacniacz operacyjny jest jednym z najlepiej opisanych elementów elektroniki, a mimo to samowzbudne oscylacje pozostają jedną z najczęstszych przyczyn, dla których prototyp nie działa tak, jak przewiduje schemat. Jak ujął to Barry Harvey w nocie Linear Technology AN148, projektanci dokładają starań, by ich wzmacniacze były stabilne, lecz w realnym układzie wiele sytuacji wytrąca je z równowagi: różne rodzaje obciążenia, źle zaprojektowana sieć sprzężenia, niewystarczające odsprężanie zasilania, a nawet samo wejście lub wyjście traktowane jako dwójnik. Większość tych zjawisk jest przewidywalna i policzalna, zanim dotknie się lutownicy. W tym artykule rozprawiamy się z ośmioma powszechnymi przekonaniem na temat stabilności wzmacniaczy operacyjnych, opierając się wyłącznie na notach aplikacyjnych i przewodnikach projektowych oraz na publikacjach pięciu inżynierów, którzy te układy projektowali i mierzyli zawodowo: Boba Pease'a, Jima Williamsa, Walta Junga, Tima Greena i Bruce'a Trumpa.

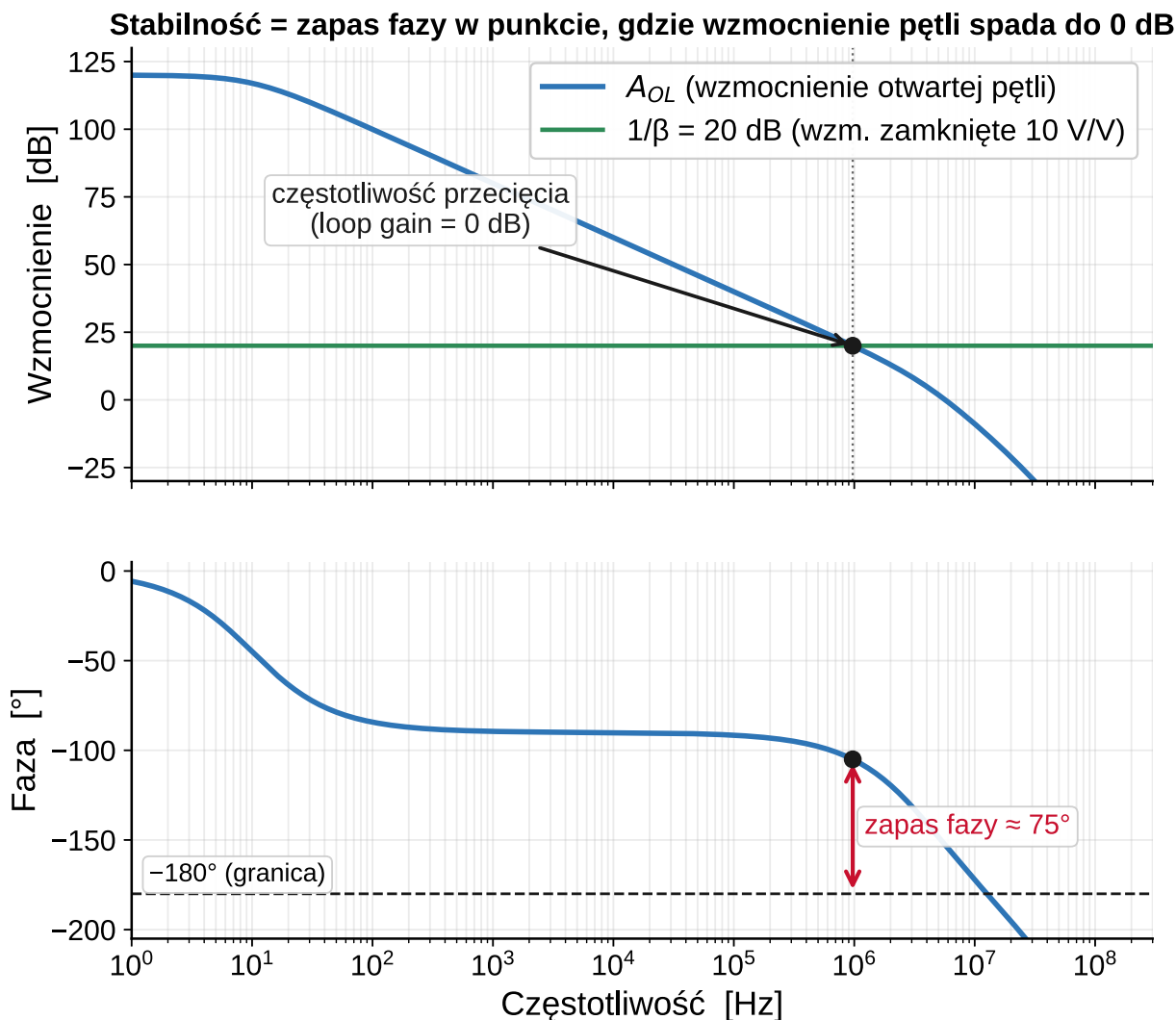
Wprowadzenie: na czym naprawdę polega stabilność

Wzmacniacz operacyjny prawie nigdy nie pracuje bez sprzężenia zwrotnego. To właśnie pętla ujemnego sprzężenia zwrotnego ustala wzmocnienie, pasmo i dokładność układu, ale ta sama pętla może stać się źródłem niestabilności. Sygnał, który wraca z wyjścia na wejście odwracające, jest opóźniony w fazie przez kolejne bieguny toru. Jeżeli przy częstotliwości, na której wzmocnienie pętli (loop gain – iloczyn wzmocnienia w otwartej pętli A_{OL} i współczynnika sprzężenia β) spada do jedności, łączne przesunięcie fazy osiągnie 360° (czyli dodatkowe 180° ponad nieodłączne 180° wejścia odwracającego), ujemne sprzężenie staje się dodatnim i układ spełnia warunek oscylacji.

Harvey ujmuje ten warunek wprost: oscylacja wystąpi, gdy opóźnienie fazy w pętli narosnie do $\pm 360^\circ$ (lub jego wielokrotności), a wzmocnienie pętli wynosi co najmniej 0 dB. Praktyczną miarą bezpieczeństwa jest **zapas fazy** (phase margin) – różnica między rzeczywistym przesunięciem fazy a granicą 360° w punkcie, w którym wzmocnienie pętli przechodzi przez 0 dB. Dla wtórniaka napięciowego, w którym na wejście wraca całe napięcie wyjściowe, warunek jest najtrudniejszy do spełnienia

– i to właśnie wtórniaki oraz układy o małym wzmocnieniu są najbardziej podatne na oscylacje. Drugą, równie ważną, choć rzadziej wspominaną miarą jest zapas wzmocnienia (gain margin): wartość wzmocnienia pętli w punkcie, w którym faza osiąga granicę. W praktyce przyjmuje się, że zdrowy układ ma zapas fazy rzędu kilkudziesięciu stopni; wartość poniżej około $35...45^\circ$ oznacza już wyraźne przeregulowanie i dzwonięcie, a zapas zerowy – oscylacje.

Skąd bierze się dodatkowe przesunięcie fazy? Z biegunów. Każdy biegun w torze sprzężenia dodaje do -90° opóźnienia powyżej swojej częstotliwości i 45° dokładnie w jej punkcie. Wewnętrznie skompensowany wzmacniacz ma celowo wprowadzony biegun dominujący, który sprowadza wzmocnienie do jedności, zanim kolejne bieguny zdążą dołożyć krytyczne 180° . Problem zaczyna się wtedy, gdy do toru dochodzi drugi biegun blisko częstotliwości przecięcia – z obciążenia, z sieci sprzężenia albo z pasożytnów płytki. Toshiba w przewodniku „Basics of Operational Amplifiers and Comparators” wyprowadza ten sam warunek od strony filtru RC: pojedynczy biegun daje 45° opóźnienia w swojej częstotliwości i dąży do 90° powyżej niej, więc dwa bieguny w okolicy przecięcia wystarczą, by zbliżyć się do granicy.



Rysunek 1. Stabilność rozstrzyga się w punkcie, w którym wzmacnienie pętli spada do 0 dB: liczy się, jak daleko faza jest wtedy od granicy -180° . Zapas fazy $\approx 75^\circ$ oznacza układ stabilny. Źródło: Barry Harvey, „Does Your Op Amp Oscillate?”, Linear Technology AN148, fig. 2 – <https://www.analog.com/media/en/technical-documentation/application-notes/an148fa.pdf>

Mit 1: „ujemne sprzężenie zwrotne samo zapewnia stabilność”

Najtrwalsze przekonanie głosi, że skoro pętla ujemnego sprzężenia koryguje błędy, to im głębsze sprzężenie, tym lepiej, a stabilność „załatwia się sama”. W rzeczywistości głębokość sprzężenia nie gwarantuje stabilności – decyduje o niej wzajemne położenie krzywej A_{OL} i krzywej odwrotności sprzężenia $1/\beta$. Tim Green w obszernym cyklu „Operational Amplifier Stability” sprowadza analizę do prostej, graficznej reguły: o stabilności rozstrzyga różnica nachyleń obu krzywych w punkcie ich przecięcia (*rate of closure*). Jeżeli krzywe zbliżają się z różnicą 20 dB/dekadę, układ ma duży zapas fazy i jest stabilny; jeżeli następuje to z różnicą 40 dB/dekadę, zapas fazy znika.

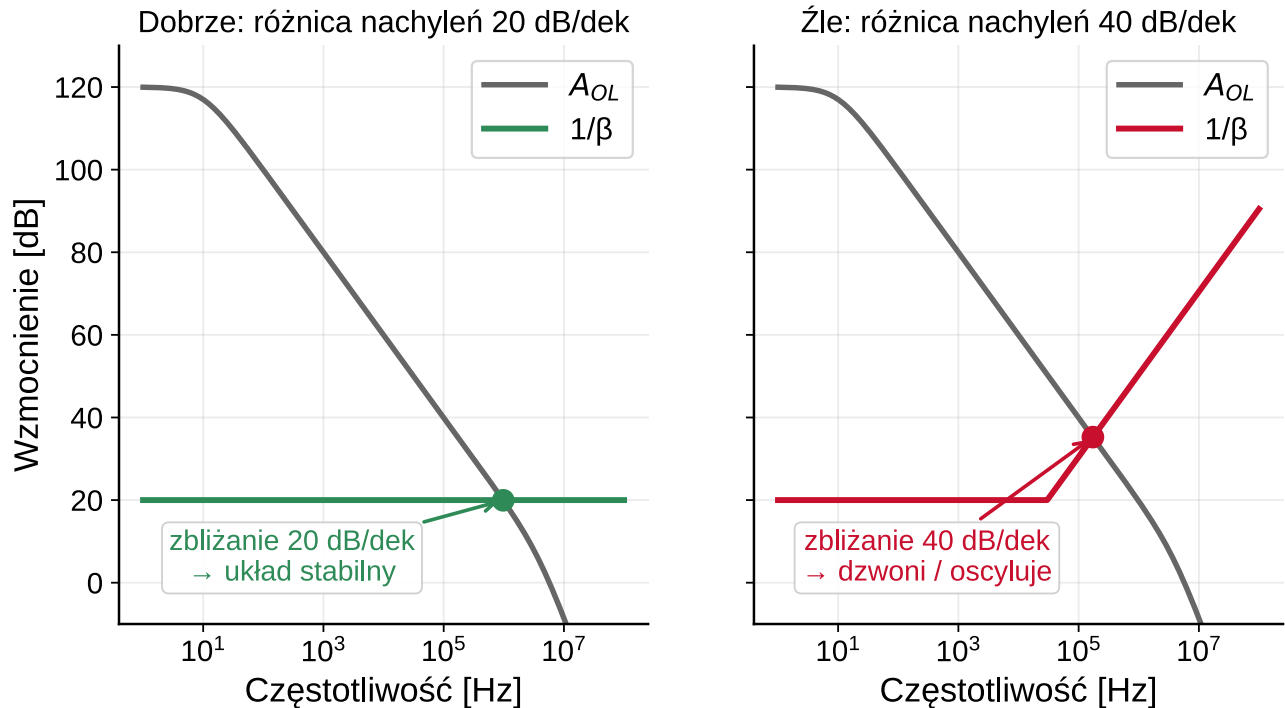
Mechanizm jest następujący. Nachylenie 20 dB/dekadę odpowiada jednemu biegunowi i co najwyżej 90° opóźnienia fazy – bezpiecznie daleko od granicy. Nachylenie 40 dB/dekadę to dwa bieguny i nawet 180° opóźnienia; jeżeli wypada ono w okolicy przecięcia, zapas fazy spada do zera. Dlatego sam fakt, że pętla jest „mocna”, nic nie

mówi o stabilności: liczy się, z jakim nachyleniem A_{OL} dochodzi do linii $1/\beta$. Ten sam mechanizm opisuje TI w przewodniku „Op Amps for Everyone” (rozdział o teorii sprzężenia i stabilności): wykresy wzmacnienia pętli są kluczem do zrozumienia stabilności, a równanie drugiego rzędu wiąże zapas fazy z przeregulowaniem i dzwonieniem odpowiedzi skokowej.

Wniosek praktyczny jest taki, że stabilności nie da się ocenić, patrząc tylko na wzmacnienie w paśmie użytecznym. Trzeba spojrzeć na całą charakterystykę A_{OL} aż do częstotliwości przecięcia z $1/\beta$ i sprawdzić, ile biegunów się tam spotyka. Kolejne mity pokazują, skąd najczęściej bierze się ten groźny drugi biegun: z obciążenia pojemnościowego, z kabla, z pojemności wejścia, z pasywnych elementów płytki.

Mit 2: „obciążenie pojemnościowe to drobiazg”

Dołączenie pojemności do wyjścia wydaje się niewinne – w końcu wzmacniacz „ma niską impedancję wyjściową”.

Reguła szybkości zbliżania (rate of closure) krzywych A_{OL} i $1/\beta$ 

Rysunek 2. Reguła szybkości zbliżania: o stabilności decyduje różnica nachyleń krzywych A_{OL} i $1/\beta$ w punkcie przecięcia, a nie samo wzmocnienie. Źródło: Tim Green, „Operational Amplifier Stability” (cykl)/TI „Op Amps for Everyone”, SLOD006, rozdz. 5 – <https://www.ti.com/lit/an/slod006b/slod006b.pdf>

Tyle że ta impedancja jest niska tylko dzięki sprzężeniu i tylko w paśmie pętli; w otwartej pętli wyjście ma realną rezystancję R_o rzędu od kilku do kilkudziesięciu omów. Razem z pojemnością obciążenia C_L tworzy ona człon dolnoprzepustowy o biegunie:

$$f_p = \frac{1}{2\pi R_o C_L}$$

ROHM w nocie „Oscillation of Op-Amp Caused by Capacitive Load” wyprowadza ten biegun wprost z transmitancji wtórnika i pokazuje, że dla danego R_o rosnące C_L przesuną biegun w dół częstotliwości – coraz głębiej w pasmo pętli – i coraz mocniej zjada zapas fazy. Wtórnik napięciowy jest tu najbardziej narażony, bo pracuje przy $1/\beta = 0$ dB, więc biegun obciążenia spotyka się z krzywą A_{OL} przy pełnym wzmocnieniu pętli. Dla $R_o = 100 \Omega$ i $C_L = 1$ nF biegun wypada przy ok. 1,6 MHz – typowo wewnątrz pasma pętli szybkiego wzmacniacza.

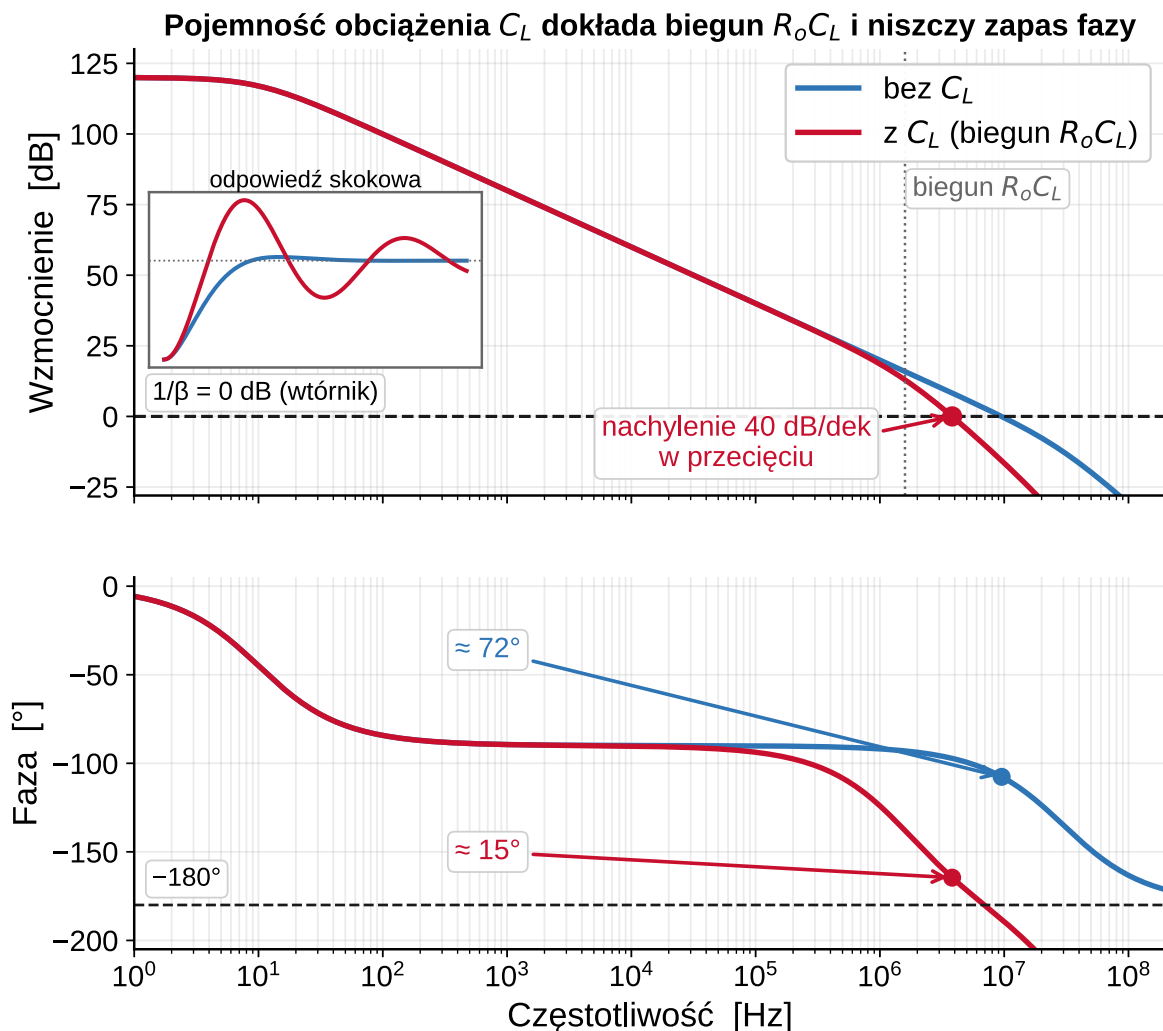
Harvey w AN148 dodaje, że problem nie kończy się na prostym biegunie. Powyżej pewnej częstotliwości impedancja wyjściowa wielu wzmacniaczy staje się indukcyjna (efekt skończonego f_T tranzystorów wyjściowych, do tego indukcyjność obudowy i wyprowadzeń, łącznie kilka–kilkadziesiąt nH). Ta indukcyjność L_{OUT} razem z C_L tworzy szeregowy obwód rezonansowy, którego impedancja w rezonansie spada do bardzo małej wartości – wyjście „widzi” niemal zwarcie i dokłada kolejne opóźnienie fazy. To dlatego niektóre układy oscylują dopiero

przy konkretnej wartości pojemności, a nie monotonicznie z jej wzrostem.

Jak to opanować: R_{iso} i podwójne sprzężenie zwrotne

Najprostszym środkiem jest rezystor izolujący R_{iso} włączony szeregowo między wyjście a pojemność. Dla wysokich częstotliwości to R_{iso} , a nie sama R_o , ustala impedancję widzianą przez C_L , co odsuwa biegun obciążenia z krytycznego obszaru pętli. Wadą jest błąd stałoprądowy: jeżeli sprzężenie nadal pobieramy sprzed R_{iso} , napięcie na obciążeniu spada na dzielniku $R_{iso}-R_{obc}$. Rozwiązaniem zachowującym dokładność jest podwójne sprzężenie zwrotne (dual feedback): składową stałoprądową pobieramy przez rezystor R_f z węzła obciążenia (za R_{iso}), a składową wielkoczęstotliwościową przez kondensator C_f wprost z wyjścia wzmacniacza (sprzed R_{iso}). Pętla „widzi” rzeczywiste napięcie na obciążeniu (dokładność DC), a dla wysokich częstotliwości omija biegun obciążenia (stabilność).

Williams w nocie AN47 „High Speed Amplifier Techniques” pokazuje cały katalog takich technik dla wzmacniaczy szybkich oraz – co równie ważne – przestrzega, że przy tych prędkościach o stabilności decyduje nie tylko schemat, lecz montaż: prowadzenie masy, odsprężanie i pasożyty płytki. Do tego wątku wracamy w micie 6.

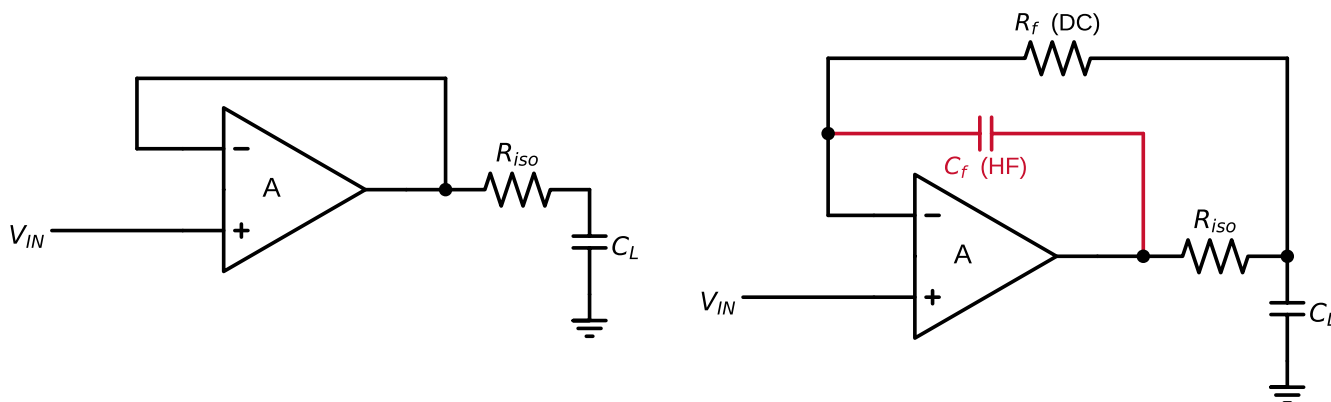


Rysunek 3. Pojemność obciążenia C_L wraz z rezystancją wyjścia R_o tworzy dodatkowy biegun w pętli; nachylenie rośnie do 40 dB/dek, faza zbliża się do -180° , a odpowiedź skokowa dzwoni. Źródło: ROHM, „Oscillation of Op-Amp Caused by Capacitive Load”, rysunek 9 i charakterystyki fazy w funkcji C_L
https://fscdn.rohm.com/en/products/databook/applinote/ic/amp_linear/opamp/gpl_opa_osc_load_capa-e.pdf

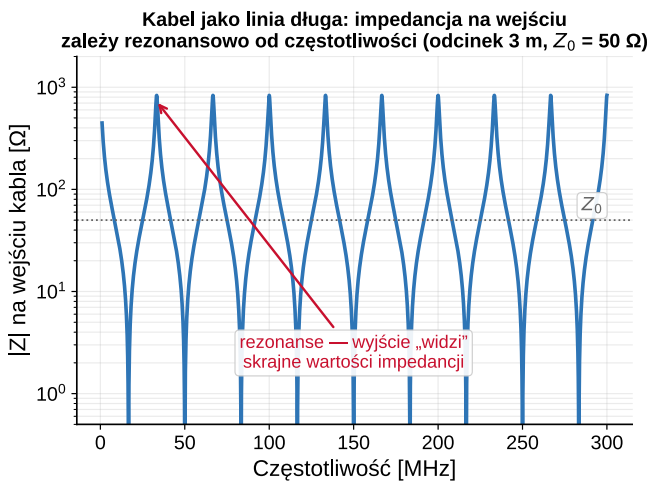
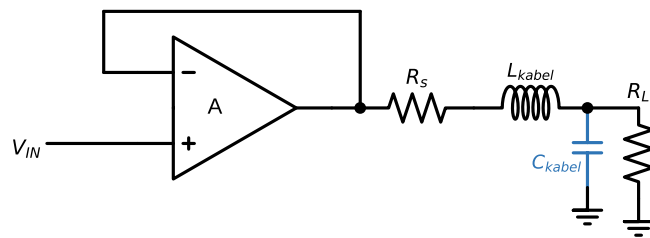
Mit 3: „kabel to tylko połączenie”

Kilka metrów ekranowanego przewodu na wyjściu wydaje się neutralne, lecz dla wzmacniacza kabel jest obciążeniem. Po pierwsze, jego pojemność (rzędu 100 pF na metr) dołącza się wprost do wyjścia

i sumuje z pojemnością obciążenia – wraca więc mit 2: rezystancja wyjścia R_o i pojemność kabla tworzą biegun, który zjada zapas fazy. Po drugie, kabel rozwarły (niezakończony) lub źle dopasowany zachowuje się jak linia długa: jego impedancja widziana od strony



Rysunek 4. Stabilizacja obciążenia pojemnościowego: (a) R_{iso} poza pętlą – stabilny, lecz wnosi błąd DC na dzielniku; (b) R_{iso} z podwójnym sprzężeniem – R_f z węzła obciążenia daje dokładność DC, C_f z wyjścia zapewnia stabilność HF. Źródło: Tim Green, „Operational Amplifier Stability” – metody R_{iso} oraz R_{iso} z podwójnym sprzężeniem; AN148 fig. 9 (snubber jako wariant)
<https://www.ti.com/lit/an/sboa015/sboa015.pdf>

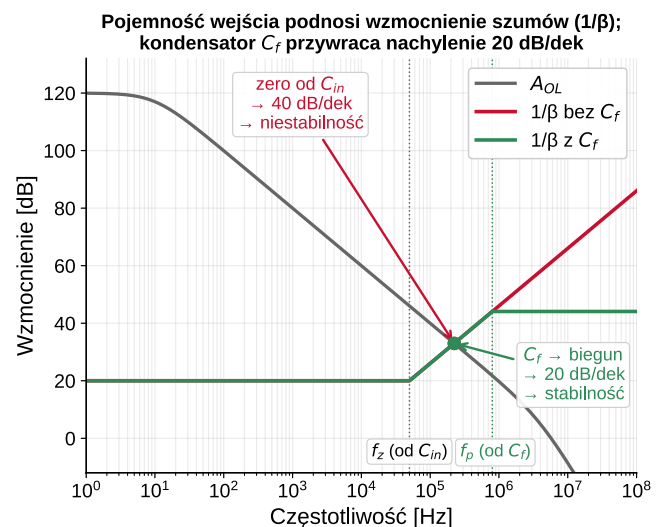
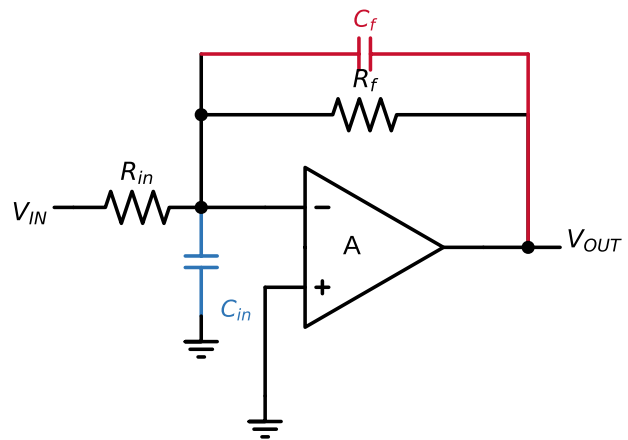


Rysunek 5. Kabel na wyjściu to obciążenie: jego pojemność C_{kabel} dokładnie biegun, a linia rozwarta (niezakończona) ma impedancję zmieniającą się rezonansowo z częstotliwością. Rezystor szeregowy R_s (back-termination) tłumi te efekty. Źródło: Barry Harvey, „Does Your Op Amp Oscillate?”, Linear Technology AN148 (impedancja kabla na wyjściu wzmacniacza) – <https://www.analog.com/media/en/technical-documentation/application-notes/an148fa.pdf>

wzmacniacza zmienia się z częstotliwością rezonansowo, a w okolicach rezonansów (długość bliska wielokrotności ćwierci fali) wyjście „widzi” skrajne wartości impedancji.

Impedancja widziana na wejściu takiego kabla nie jest stała – zmienia się z częstotliwością i w okolicach rezonansów osiąga wartości skrajne (rysunek 5). Indukcyjność i pojemność samego przewodu, oznaczone na naszym schemacie L_{kabel} i C_{kabel} , mogą przy pewnych częstotliwościach utworzyć obwód rezonansowy, w którym wyjście wzmacniacza obciążone jest bardzo małą impedancją, co dokłada kolejne opóźnienie fazy w pętli. To jeden z mechanizmów, dla których układ stabilny na stole zaczyna oscylować dopiero po dołączeniu konkretnego odcinka kabla.

Williams w nocie AN47, w „galerii problemów” („Mr. Murphy’s Gallery”), umieszcza źle zakończoną linię wśród typowych przyczyn zniekształceń i niestabilności szybkich torów. Praktyczne środki są dwa: traktować kabel jako część obciążenia (i stabilizować jak w micie 2 – rezystorem szeregowym lub podwójnym sprzężeniem), a gdy przewód prowadzi sygnał na większą odległość, zakończyć linię, na przykład rezystorem szeregowym o wartości bliskiej impedancji falowej kabla (back-termination), co tłumi odbicia i wygładza impedancję widzianą przez wzmacniacz.



Rysunek 6. Pojemność C_{in} na wejściu odwracającym tworzy zero w $1/\beta$ (wzmocnieniu szumów): nachylenie rośnie do 40 dB/dek i układ traci stabilność. Kondensator C_f równoległe do R_f dokłada biegun i przywraca nachylenie 20 dB/dek. Źródło: Barry Harvey, AN148 (metody ograniczania wpływu pojemności wejścia)/TI „Op Amps for Everyone”, SLOD006, rozdz. 8 – <https://www.analog.com/media/en/technical-documentation/application-notes/an148fa.pdf>

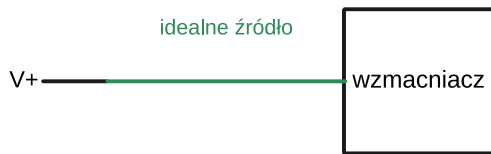
Mit 4: „pojemność na wejściu odwracającym jest nieszkodliwa”

Na węzle sumującym (wejściu odwracającym) zawsze jest jakaś pojemność: pojemność wejściowa samego wzmacniacza, pojemność montażowa ścieżek, pojemność elementów. Harvey zauważa, że każdy węzeł na płycie ma typowo co najmniej 2 pF, a wzmacniacze niskoszumne i o małym napięciu niezrównoważenia mają duże tranzystory wejściowe, więc i większą pojemność wejścia. Ta pojemność C_{in} wraz z rezystancją sieci sprzężenia tworzy zero w przebiegu odwrotności sprzężenia $1/\beta$ (czyli we wzmocnieniu szumów). Zero podnosi $1/\beta$ z nachyleniem 20 dB/dekadę, więc w punkcie przecięcia z A_{OL} różnica nachyleń rośnie do 40 dB/dekadę – i wracamy do warunku z mitu 1: zapas fazy znika.

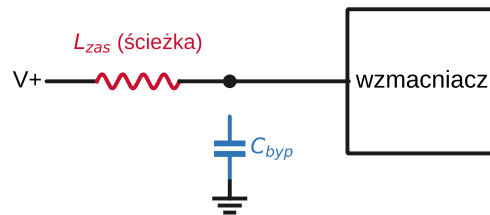
Częstotliwość zera jest wyznaczona przez pojemność wejścia i rezystancję obwodu sprzężenia (w przybliżeniu – dla równoległego połączenia rezystorów obwodu):

Mit „działa w SPICE = zadziała na płytce”: czego nie ma w modelu

Model w SPICE — szyna idealna

brak L_{zas} , brak C_{byp}

Rzeczywista płytko — szyna realna



indukcyjność ścieżki + odsprężanie tuż przy pinie

(dla szyny V- analogicznie). Green w cz. 7 cyklu pokazuje, że nawet samo R_O z modelu SPICE potrafi różnić się od pomiaru ok. pięciokrotnie.

Rysunek 7. Czego nie ma w modelu: idealne szyny zasilania w SPICE nie mają indukcyjności ścieżek ani potrzeby odsprężania; na płytce L_{zas} i C_{byp} są realne. Źródło: Jim Williams, „High Speed Amplifier Techniques”, AN47 (odsprężanie i montaż prototypu) – <https://www.analog.com/media/en/technical-documentation/application-notes/an47fa.pdf>

$$f_z \approx \frac{1}{2\pi R_{eq} C_{in}}$$

Lekarstwo jest dobrze znane: kondensator C_f równoległy do rezystora sprzężenia R_f . Wprowadza on biegun, który spłaszcza $1/\beta$ z powrotem do nachylenia 20 dB/dekadę przed punktem przecięcia:

$$f_p \approx \frac{1}{2\pi R_f C_f}$$

Harvey podaje też prostą regułę doboru: biegun sprzężenia powinien leżeć co najmniej sześciokrotnie poniżej częstotliwości jednostkowego wzmocnienia, a wartość rezystora sprzężenia trzeba dobierać z uwzględnieniem pojemności pasożytniczej, by biegun ten nie wypadł zbyt blisko częstotliwości przecięcia. To samo zagadnienie – stabilność a pojemność wejścia i sprzężenia – TI omawia w „Op Amps for Everyone” (rozdział o kompensacji), a Tim Green ujmuje je w kategoriach kształtowania wzmocnienia szumów. Uwaga praktyczna: problem dotyczy przede wszystkim wzmacniaczy ze sprzężeniem napięciowym; wzmacniacze ze sprzężeniem prądowym nie

tolerują kondensatora wprost w sprzężeniu i wymagają innego podejścia.

Mit 5: „skoro działa w SPICE, zadziała na płytce”

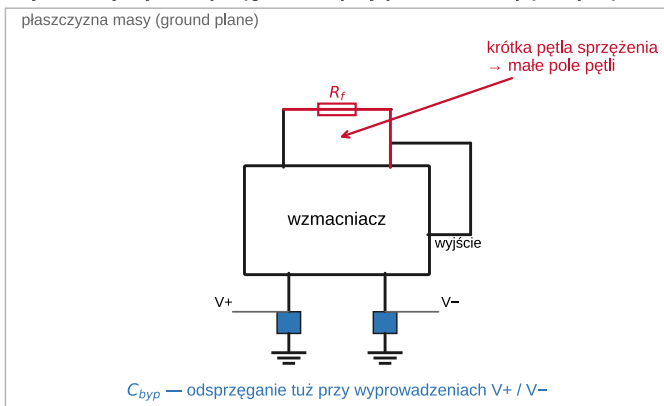
Symulacja jest użytecznym narzędziem, lecz jej wynik zależy od tego, co zawiera model. Standardowy makromodel wzmacniacza często nie odwzorowuje dokładnie przebiegu impedancji wyjściowej w funkcji częstotliwości, bywa, że pomija drugorzędne bieguny stopnia wyjściowego, a niemal nigdy nie zawiera elementów pasożytniczych płytki: indukcyjności ścieżek zasilania, pojemności montażowych i sprzężeń. Idealne szyny zasilania w symulacji nie wymagają odsprężania – na płytce brak kondensatorów odsprężających tuż przy wyprowadzeniach potrafi zamienić poprawny układ w generator.

Najlepiej ilustruje to sam Green: w 7. części cyklu o stabilności zmierzył rezystancję wyjściową R_O rzeczywistego wzmacniacza CMOS z wyjściem rail-to-rail i porównał ją z symulacją TINA-SPICE – wynik modelu był zaniżony około pięciokrotnie. Ponieważ to właśnie R_O decyduje o położeniu bieguna obciążenia (mit 2), błąd o czynnik pięć oznacza całkowicie błędną ocenę zapasu fazy. Williams w AN47 podsumowuje tę postawę krótko: niczego nie należy przyjmować z góry. Wniosek praktyczny: kolejność to analiza, symulacja, a potem bezwzględnie pomiar na prototypie – z prawdziwym obciążeniem, kablem i odsprężaniem.

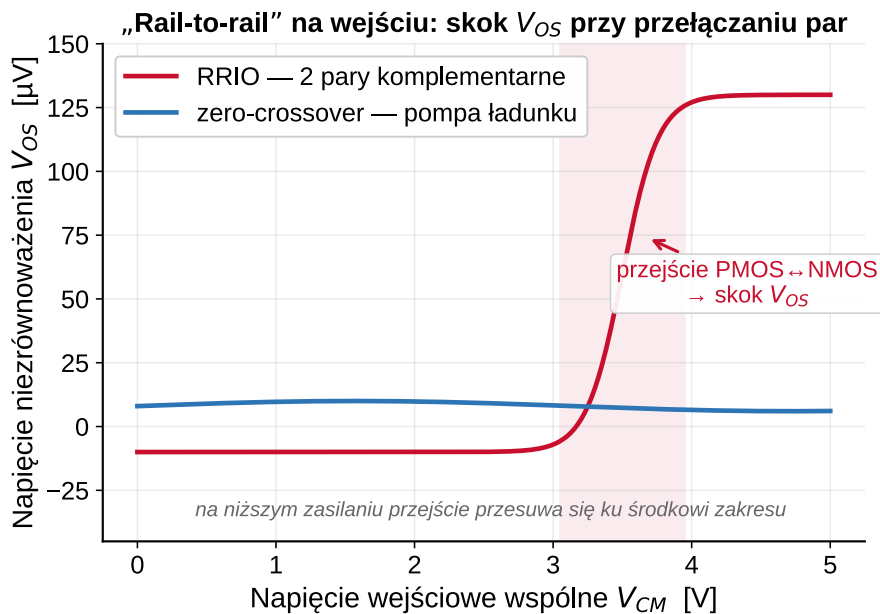
Mit 6: „layout to kosmetyka”

Schemat ideowy nie pokazuje indukcyjności ścieżek, pól pętli ani sprzężeń pojemnościowych – a to one przy wyższych częstotliwościach decydują o stabilności. Williams poświęcił temu w AN47 obszerną „galerię problemów”, w której powtarzają się te same przyczyny: brak kondensatorów odsprężających przy dużym obciążeniu, brak płaszczyzny masy oraz sprzężenia pojemnościowe

Layout decyduje: odsprężanie tuż przy pinach i mała pętla sprzężenia



Rysunek 8. Trzy zasady z praktyki: kondensatory odsprężające C_{byp} tuż przy wyprowadzeniach zasilania, ciągła płaszczyzna masy oraz możliwie małe pole pętli sprzężenia R_f . Źródło: Jim Williams, „High Speed Amplifier Techniques”, AN47 (montaż prototypu, „Mr. Murphy’s Gallery”) – <https://www.analog.com/media/en/technical-documentation/application-notes/an47fa.pdf>



Rysunek 9. Wejście RRIO z dwiema komplementarnymi parami wykazuje skok V_{OS} w obszarze przejścia PMOS ↔ NMOS; topologia zero-crossover (pompa ładunku) utrzymuje płaski przebieg. Na niższym zasilaniu obszar przejścia przesuwa się ku środkowi zakresu. Źródło: Bruce Trump, „The Signal” – „Rail-to-Rail Inputs—what you should know!”, Texas Instruments E2E – https://e2e.ti.com/blogs_/archives/b/thesignal/posts/rail-to-rail-inputs-what-you-should-know

do krytycznych, wysokoomowych węzłów (np. węzła sumującego). Każda z nich potrafi samodzielnie wywołać oscylacje w układzie, który „na papierze” jest poprawny.

Środki zaradcze są proste i tanie, o ile zastosuje się je od początku projektu płytki. Kondensator odsprężający należy umieścić jak najbliżej wyprowadzenia zasilania, by zminimalizować indukcyjność doprowadzenia, przez którą płyną szybkie prądy stopnia wyjściowego. Ciągła płaszczyna masy skraca drogę powrotu prądu i obniża impedancję odniesienia. Pętlę sprzężenia – zwłaszcza odcinek do wejścia odwracającego – trzeba prowadzić krótko, aby zmniejszyć pole pętli i podatność na sprzężenia. TI poświęca tym zagadnieniom osobny rozdział w „Op Amps for Everyone”, traktując layout nie jako kosmetykę, lecz jako element obwodu. Do mechanizmu zasilania i odsprężania wracaliśmy już w micie 5 – to dwie strony tego samego zagadnienia.

Mit 7: „wzmacniacz rail-to-rail jest zawsze lepszy”

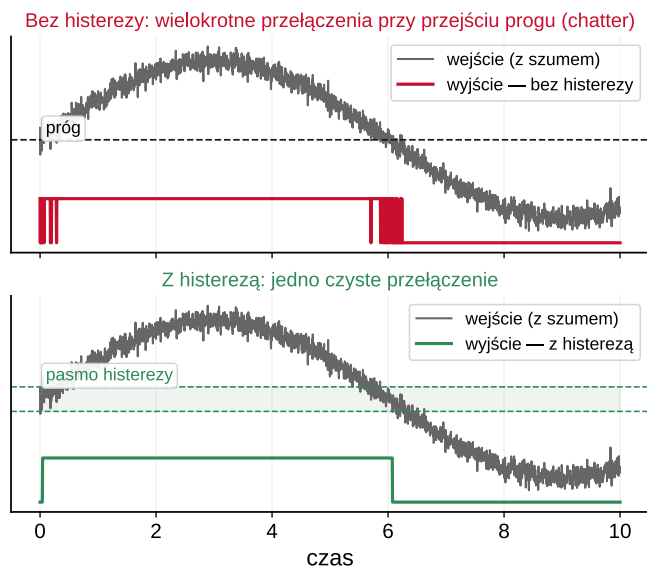
Wejście „rail-to-rail” (RRIO) kusi obietnicą wykorzystania całego zakresu zasilania, lecz w typowej realizacji ma cenę, o której rzadko mówi się na wykładach. Aby objąć napięcie wspólne od jednej szyny do drugiej, producenci łączą zwykle dwie komplementarne pary wejściowe – parę PMOS dla napięć bliskich ujemnej szyny i parę NMOS dla napięć bliskich dodatniej. Każda z par ma własne napięcie niezrównoważenia, więc w obszarze, w którym sterowanie przechodzi z jednej pary na drugą, napięcie niezrównoważenia V_{OS} wykonuje skok. Bruce Trump w cyklu „The Signal” opisuje to wprost: zmiana V_{OS} przy przełączaniu par jest źródłem zniekształceń (crossover distortion) w zastosowaniach zmiennoprądowych i błędów w precyzyjnych zastosowaniach stałoprądowych.

Co istotne, układ sterujący przejściem między parami jest odniesiony do dodatniej szyny zasilania, a nie do masy. Na zasilaniu 5 V obszar przejścia leży blisko górnej szyny i często pozostaje poza użytecznym zakresem sygnału; ale przy zasilaniu 3,3 V – jak zauważa Trump – przesuwa się ku środkowi zakresu, czyli w najgorsze możliwe miejsce. Fabryczne dostrajanie (laserowe lub elektroniczne) zmniejsza skok, lecz pozostawia resztkowe „drżenie” (bobble). Rozwiązaniem, które całkowicie eliminuje przejście, jest topologia zero-crossover: pojedyncza para wejściowa zasilana z wewnętrznej pompy ładunku, która podnosi jej napięcie zasilania o ok. 1,5...2 V powyżej dodatniej szyny, dzięki czemu jedna para obsługuje cały zakres bez przełączania.

Praktyczny wniosek nie brzmi „nigdy nie używać RRIO”, lecz „używać świadomie”. W układzie odwracającym lub transimpedancyjnym napięcie wspólne wejść jest ustalone i nie przemierza obszaru przejścia, więc problem znika. Gdy jednak napięcie wspólne podąża za sygnałem (np. wtórnik, wzmacniacz nieodwracający), trzeba sprawdzić w karcie katalogowej, gdzie leży obszar przejścia i jaki jest rozrzut V_{OS} . Warto też pamiętać, że osobnym zagadnieniem jest zakres napięcia wspólnego wejść w układach jednonapięciowych: nie każdy wzmacniacz obejmuje masę, a założenie, że „wejście sięga do zera”, bywa źródłem błędów równie częstym jak sama oscylacja.

Mit 8: „wzmacniacz operacyjny zadziała jako komparator”

Pokusa jest zrozumiała: skoro wzmacniacz operacyjny ma ogromne wzmocnienie, to porówna dwa napięcia równie dobrze jak komparator. W praktyce użycie wzmacniacza operacyjnego jako komparatora prowadzi do trzech kłopotów. Po pierwsze, wzmacniacze operacyjne



Rysunek 10. Powolne przejście sygnału przez próg w obecności szumu: bez histerezy wyjście przełącza się wielokrotnie (chatter), z histerezą (pasmo dwóch progów) następuje jedno czyste przełączenie. Źródło: Toshiba, „Basics of Operational Amplifiers and Comparators” (komparatory, histereza)/J. Williams, AN47 – <https://www.analog.com/media/en/technical-documentation/application-notes/an47fa.pdf>

mają wewnętrzny kondensator kompensujący, który celowo spowalnia ich reakcję, by zapewnić stabilność w pętli; w roli komparatora oznacza to powolne wychodzenie z nasycenia i rozmyte zbocza. Toshiba w przewodniku „Basics of Operational Amplifiers and Comparators” zaznacza, że komparatory świadomie nie mają takiej kompensacji – i właśnie dlatego są szybkie.

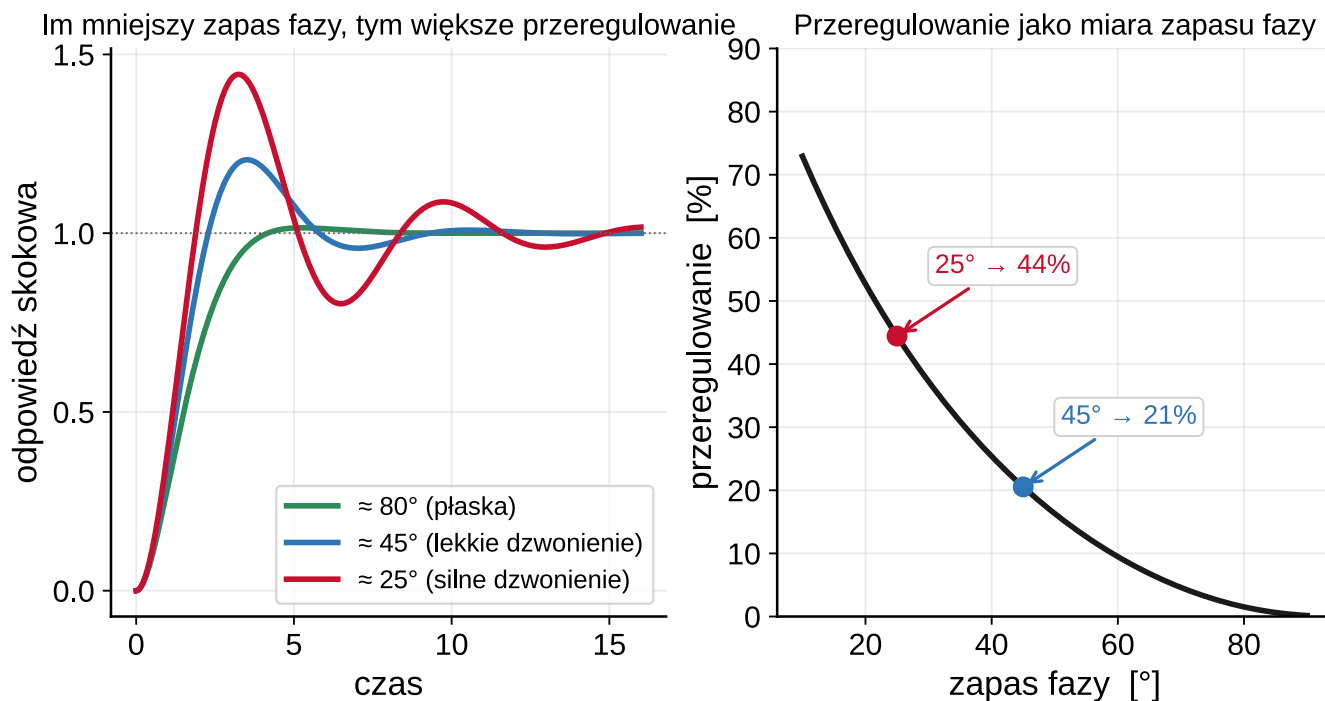
Po drugie, wzmacniacz operacyjny pracujący bez pętli przy powolnym przejściu sygnału przez próg jest

podatny na wielokrotne przełączenia (chatter), gdy do sygnału dołoży się szum. Lekarstwem jest histereza – niewielkie dodatnie sprzężenie zwrotne, które rozsuwa próg załączania i wyłączania, dając jedno czyste przełączenie zamiast serii drgań. Toshiba formułuje tu jednoznaczną regułę: nie należy stosować pętli sprzężenia wokół komparatora – z wyjątkiem właśnie histerezy (dodatniego sprzężenia). Po trzecie, przekroczenie dopuszczalnego zakresu napięcia wspólnego wejść grozi w niektórych konstrukcjach odwróceniem fazy (phase reversal), czyli nagłą zmianą znaku wyjścia – co w pętli regulacji bywa katastrofalne.

Wniosek: do porównywania napięć należy użyć komparatora (z histerezą dobraną do poziomu szumu), a wzmacniacz operacyjny zostawić do pracy liniowej w pętli ujemnego sprzężenia. Jeżeli z jakichś powodów trzeba użyć wzmacniacza operacyjnego jako komparatora, histereza jest obowiązkowa, a kartę katalogową trzeba sprawdzić pod kątem zachowania przy przekroczeniu zakresu wejść.

Warsztat A: jak zmierzyć zapas fazy i wykryć oscylacje

Skoro o stabilności decyduje zapas fazy, przydaje się sposób na jego oszacowanie bez analizatora. Najprostszy opiera się na odpowiedzi skokowej. Układ drugiego rzędu wiąże zapas fazy z przeregulowaniem: im mniejszy zapas, tym większe przeregulowanie i dłuższe dzwonienie. „Op Amps for Everyone” podaje tę zależność wprost – odpowiedź na skok pozwala odczytać względną stabilność. W praktyce: podajemy na wejście mały skok napięcia



Rysunek 11. Po lewej: odpowiedź skokowa dla malejącego zapasu fazy (rosnące przeregulowanie i dzwonienie). Po prawej: przeregulowanie jako miara zapasu fazy – z odczytu na oscyloskopie można oszacować margines. Źródło: TI „Op Amps for Everyone”, SLOD006 (układ drugiego rzędu, przeregulowanie)/TI Precision Labs – Op Amps: Stability – <https://www.ti.com/lit/an/slod006b/slod006b.pdf>

i obserwujemy wyjście. Brak przeregulowania to zapas rzędu 70...80°; przeregulowanie kilkunastu procent odpowiada ok. 45°; przeregulowanie rzędu 40% i wyraźne dzwonięcie to już ok. 25° – układ na granicy.

Trzy sposoby wykrycia oscylacji na stanowisku

Gdy podejrzewamy samowzbudzenie, materiały TI („Troubleshooting Tips: Op Amps – Oscillations”) wskazują trzy proste metody. Pierwsza: pomiar prądu zasilania – oscylujący stopień wyjściowy pobiera zauważalnie większy prąd średni niż w spoczynku, więc wzrost poboru bywa pierwszym sygnałem. Druga: obserwacja wyjścia oscyloskopem z odpowiednim pasmem – przy czym trzeba pamiętać o ostrzeżeniu Williamsa z AN47, że źle uziemiona lub przeciążona sonda sama potrafi „pokazać” albo zamaskować oscylacje. Trzecia: użycie małej cewki lub odcinka przewodu jako sondy pola bliskiego („anten”) trzymanej nad płytką – pozwala zlokalizować źródło promieniowania wielkiej częstotliwości bez galwanicznego dotaczania się do wrażliwego węzła.

Warsztat B: sześć sposobów na stabilne obciążenie pojemnościowe

Tim Green w cyklu „Operational Amplifier Stability” zebrał sześć technik stabilizacji wzmacniacza obciążonego pojemnością. Dwie z nich poznaliśmy już przy micie 2 (rysunek 4). Poniższe zestawienie porządkuje wszystkie sześć wraz z zasadą działania – wybór zależy od tego, czy ważniejsza jest dokładność stałoprądowa, pasmo, czy prostota układu.

Metoda	Zasada działania/kiedy stosować
Rezystor izolujący R_{iso}	Szeregowy rezystor odsuwa biegun obciążenia; najprostsza metoda, lecz wnosi błąd stałoprądowy na obciążeniu.
Duże wzmocnienie + C_f	Praca przy większym wzmocnieniu zamkniętym plus kondensator sprzężenia; układ stabilny kosztem braku pracy przy wzmocnieniu 1.
Wzmocnienie szumów (Noise Gain)	Podniesienie wzmocnienia szumów w wielkiej częstotliwości, bez zmiany wzmocnienia sygnału; poprawia margines.
Wzmocnienie szumów + C_f	Połączenie poprzedniej metody z kondensatorem sprzężenia, ograniczające wzrost szumu w paśmie.
Kompensacja na pinie wyjściowym	Sieć RC dołączona do wyjścia (output pin compensation); odmienna od snubbera, kształtuje impedancję wyjścia.
R_{iso} z podwójnym sprzężeniem	R_{iso} plus dwie drogi sprzężenia (DC z obciążenia, HF z wyjścia); zachowuje dokładność i stabilność – rysunek 4.

Źródło zestawienia: Tim Green, „Operational Amplifier Stability” (cykl 15 części)
– <http://educyclopedia.karadimov.info/library/acqt0704.pdf>

W skrócie: wzmacniacz operacyjny a komparator

Cecha	Wzmacniacz operacyjny	Komparator
Kompensacja częstotliwościowa	Tak – wewnętrzny kondensator	Nie – celowo brak
Reakcja na nasycenie	Powolna (rozmyte zbocza)	Szybka
Pętla sprzężenia	Ujemna – do pracy liniowej	Tylko histereza (dodatnia)
Wyjście	Liniowe, analogowe	Dwustanowe, logiczne
Podatność na oscylacje	Wtórnik/małe wzm. – wysoka	Bez histerezy – chatter
Źródło: Toshiba, „Basics of Operational Amplifiers and Comparators”; J. Williams, AN47		

Lista kontrolna stabilności (przed uruchomieniem)

- Sprawdź szybkość zbliżania krzywych A_{OL} i $1/\beta$: różnica 20 dB/dek = stabilnie, 40 dB/dek = kłopot.
- Policz biegun $R_o \cdot C_L$ dla rzeczywistej pojemności obciążenia; jeśli wpada w pasmo pętli – zastosuj R_{iso} lub podwójne sprzężenie.
- Potraktuj kabel na wyjściu jako obciążenie (pojemność + linia długa); rozważ rezystor szeregowy (back-termination).
- Oszacuj pojemność węzła sumującego C_{in} i dodaj kondensator sprzężenia C_f , jeśli $1/\beta$ rośnie ku przecięciu.
- Nie ufaj samej symulacji: zweryfikuj R_{O} i stabilność pomiarem na prototypie z realnym obciążeniem.
- Umieść kondensatory odsprężające tuż przy wyprowadzeniach zasilania; zapewnij ciągłą płaszczyznę masy.
- Przy wejściu rail-to-rail sprawdź położenie obszaru przejścia V_{OS} i rozrzut na docelowym zasilaniu.
- Przy wejściu rail-to-rail sprawdź położenie obszaru przejścia V_{OS} i rozrzut na docelowym zasilaniu.
- Do porównywania napięć użyj komparatora z histerezą – nie wzmacniacza operacyjnego.
- Zmierz zapas fazy z przeregulowania odpowiedzi skokowej ($\approx 70...80^\circ$ bez przeregulowania, $\approx 45^\circ$ kilkanaście %).
- W razie podejrzenia oscylacji sprawdź pobór prądu zasilania, obejrzyj wyjście i użyj sondy pola bliskiego.
- Zweryfikuj stabilność w skrajnych warunkach (temperatura, rozrzut wzmocnienia, maksymalne C_L).
- Sprawdź, czy wzmacniacz jest stabilny przy docelowym wzmocnieniu zamkniętym (część układów nie jest stabilna przy wzmocnieniu 1).

Sylwetki przywoływanych ekspertów

Bob Pease (1940–2011) Inżynier analogowy związany przez dziesięciolecia z National Semiconductor, projektant układów takich jak przetwornik napięcie–częstotliwość

LM331 i stabilizatory napięcia. Autor klasycznej książki „Troubleshooting Analog Circuits” oraz długoletniej, cenniejszej rubryki w „Electronic Design”. Znany z nieufności wobec ślepego polegania na symulacji i z nacisku na pomiar oraz intuicję obwodową. Zginął w wypadku samochodowym w czerwcu 2011 roku, wracając z uroczystości żałobnej swojego przyjaciela Jima Williamsa.

Jim Williams (1948–2011) Projektant analogowy i autor techniczny, związany z MIT, National Semiconductor i Linear Technology. Napisał ponad 350 publikacji, w tym słynną notę aplikacyjną AN47 „High Speed Amplifier Techniques”, do dziś uznawaną za wzorzec inżynierskiego pisarstwa. Łączył dogłębną teorię z praktyką pomiarową i montażową (breadboarding). Zmarł w czerwcu 2011 roku, kilka dni po udarze.

Walt Jung Inżynier i autor, którego „IC Op-Amp Cookbook” (1974) stało się jedną z najważniejszych książek o praktycznych zastosowaniach wzmacniaczy operacyjnych. Redaktor i współautor obszernego „Op Amp Applications Handbook” wydanego przez Analog Devices. Jego prace łączą rygor z przystępnością i są wielokrotnie

cytowane w notach aplikacyjnych. Pozostaje jednym z najczęściej przywoływanych autorytetów w dziedzinie układów analogowych.

Tim Green Inżynier aplikacyjny związany z Burr-Brown, a po przejęciu – z Texas Instruments (Tucson). Autor obszernego, piętnastoczęściowego cyklu „Operational Amplifier Stability”, który usystematyzował praktyczną analizę stabilności (m.in. regułę szybkości zbliżania i sześć metod stabilizacji obciążenia pojemnościowego). Jego materiały łączą analizę teoretyczną z symulacją i pomiarem. Cykl pozostaje jednym z najczęściej polecanych wprowadzeń do tematu.

Bruce Trump Wieloletni inżynier i menedżer aplikacji w Burr-Brown, a następnie w Texas Instruments. Autor popularnego bloga „The Signal” na forum TI E2E, w którym przystępnie wyjaśniał subtelnosci wzmacniaczy operacyjnych – w tym zachowanie wejść rail-to-rail. Jego teksty są cenione za praktyczne, oparte na pomiarach podejście. Po przejściu na emeryturę pozostają często cytowanym źródłem wiedzy inżynierskiej.

WM, DS

Bibliografia

1. B. Harvey, „Does Your Op Amp Oscillate?”, Linear Technology AN148, 2014. <https://www.analog.com/media/en/technical-documentation/application-notes/an148fa.pdf>
2. J. Williams, „High Speed Amplifier Techniques”, Linear Technology AN47, 1991. <https://www.analog.com/media/en/technical-documentation/application-notes/an47fa.pdf>
3. ROHM, „Oscillation of Op-Amp Caused by Capacitive Load”, nota aplikacyjna. https://fscdn.rohm.com/en/products/databook/applinote/ic/amp_linear/opamp/gpl_opa_osc_load_capa-e.pdf
4. Toshiba, „Basics of Operational Amplifiers and Comparators”, nota aplikacyjna (TC75S67TU).
5. R. Mancini (red.), „Op Amps for Everyone”, Texas Instruments, SLOD006. <https://www.ti.com/lit/an/slod006b/slod006b.pdf>
6. T. Green, „Operational Amplifier Stability” (cykl 15 części), Burr-Brown/Texas Instruments. <http://educypedia.karadimov.info/library/acqt0704.pdf>
7. B. Trump, „Rail-to-Rail Inputs-what you should know!”, blog „The Signal”, TI E2E. https://e2e.ti.com/blogs_/archives/b/thesignal/posts/rail-to-rail-inputs-what-you-should-know
8. TI, „Troubleshooting Tips: Op Amps – Oscillations”, materiał szkoleniowy.
9. TI Precision Labs – Op Amps: Stability, seria wykładów wideo, Texas Instruments.
10. W. Jung, „Op Amp Applications Handbook”, Analog Devices/Newnes, 2005.
11. W. Jung, „IC Op-Amp Cookbook”, Howard W. Sams, 1974.
12. R. Pease, „Troubleshooting Analog Circuits”, Butterworth-Heinemann, 1991.
13. J. Williams (red.), „The Art and Science of Analog Circuit Design”, Butterworth-Heinemann, 1995.
14. S. Franco, „Design with Operational Amplifiers and Analog Integrated Circuits”, McGraw-Hill.
15. Analog Devices, karty katalogowe LTC6268/OPA388 (przykłady wejść RRIO i zero-crossover).

REKLAMA

Publikujemy dla projektantów i programistów elektroniki

ELPORTAL.pl

50 Ω

Mit 50 omów

Dlaczego „dopasowanie do 50 Ω ”
nie gwarantuje poprawnego działania
– i co naprawdę mówi wykres Smitha

„Dopasowane do 50 Ω ” to jedno z najczęściej powtarzanych zaklęć w elektronice wysokich częstotliwości – i jedno z najgorzej rozumianych. Za tą frazą kryją się cztery nieporozumienia. Po pierwsze, 50 Ω nie jest żadną fundamentalną stałą natury, lecz historycznym kompromisem. Po drugie, „dopasowanie” nie oznacza, że coś „jest 50-omowe” – oznacza brak odbicia fali na linii o określonej impedancji charakterystycznej. Po trzecie, dopasowanie na maksimum przenoszonej mocy to nie to samo co dopasowanie na minimum szumów – a myli je nawet część podręczników. Po czwarte, niski współczynnik odbicia zmierzony analizatorem nie znaczy, że „układ działa”: idealnie dopasowany tor potrafi nie promieniować, oscylować albo mieć fatalny szum. Rozbierzmy te cztery mity po kolei – do samych podstaw.

1. Skąd się wzięło 50 Ω – kompromis, nie magia

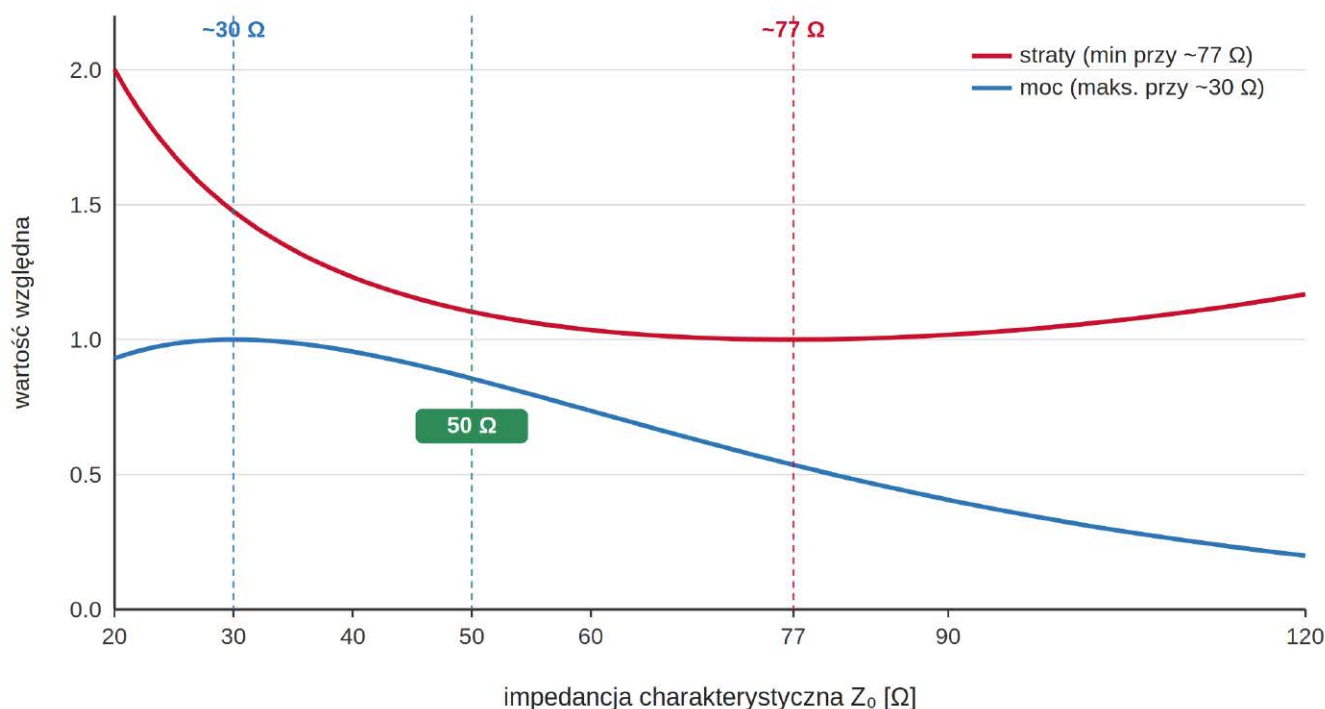
Zacznijmy od liczby, która stała się domyślną impedancją całej techniki RF (ang. *Radio Frequency*). Wartość 50 Ω nie wynika z żadnego prawa fizyki – jest wynikiem kompromisu między dwiema sprzecznymi optymalizacjami kabla współosiowego. Przyjrzyjmy się obu, bo dopiero one tłumaczą, dlaczego akurat 50.

Historia tej liczby zaczyna się w latach dwudziestych XX wieku, w Bell Telephone Laboratories. Dwaj inżynierowie, Lloyd Espenschied i Herman Affel, opracowali tam pierwszy nowoczesny kabel współosiowy – przewód, w którym jeden przewód biegnie wewnątrz drugiego, współosiowo (stąd nazwa), oddzielony dielektrykiem. Zgłoszenie patentowe z 1929 roku (patent US 1 835 031,

przyznany w 1931) rozwiązało problem, z którym nie radziła sobie telefonia dalekosiężna: taki kabel miał stabilną impedancję i małe straty przy wysokich częstotliwościach, dzięki czemu jednym łączem można było prowadzić tysiące rozmów naraz, a później także sygnał telewizyjny.

Zostało pytanie: jaką impedancję nadać nowemu kablowi? Odpowiedź kryje się w geometrii. Impedancja charakterystyczna kabla współosiowego zależy niemal wyłącznie od stosunku średnic – wewnętrznej średnicy przewodnika zewnętrznego (D) do średnicy przewodnika wewnętrznego (d) – zgodnie z zależnością $Z_0 = (60/\sqrt{\epsilon_r}) \cdot \ln(D/d)$, gdzie dla dielektryka powietrznego $\epsilon_r = 1$. Ten jeden stosunek D/d decyduje jednocześnie o tłumieniu i o wytrzymałości kabla, a każdy z tych parametrów chce innej proporcji – dlatego jednej „najlepszej” impedancji nie ma.

Skąd 50 Ω: kompromis strat i mocy w kablu współosiowym



1. Dwie sprzeczne optymalizacje kabla współosiowego. Straty mają minimum przy ok. 77 Ω, zdolność przenoszenia mocy – maksimum przy ok. 30 Ω. 50 Ω to praktyczny kompromis: blisko minimum strat, z zachowaniem ok. 85% szczytowej mocy

Bell Labs zmierzyło te optima. Najmniejsze tłumienie wypada przy $D/d \approx 3,6$ ($Z_0 \approx 77 \Omega$): zbyt cienki przewodnik wewnętrzny ma dużą rezystancję i grzeje się, a zbyt gruby niepotrzebnie zwiększa pojemność linii, więc minimum strat leży pośrodku. Największą moc kabel przenosi przy $D/d \approx 1,65$ ($Z_0 \approx 30 \Omega$): grubszy przewodnik wewnętrzny rozkłada pole elektryczne na większej powierzchni i obniża jego szczyt tuż przy przewodniku, dzięki czemu dielektryk przebija się dopiero przy wyższym napięciu i większej mocy. Największą wytrzymałość napięciową daje z kolei $D/d \approx 2,7$ ($Z_0 \approx 60 \Omega$). Trzy cele – trzy różne proporcje średnic i trzy różne liczby, leżące w różnych miejscach (rysunek 1).

50 Ω to praktyczny kompromis między mocą a tłumieniem – leży mniej więcej między 30 a 77 Ω (średnia arytmetyczna to 53,5 Ω, geometryczna ok. 48 Ω): daje straty tylko nieznacznie większe od minimalnych, a jednocześnie ok. 85% szczytowej zdolności przenoszenia mocy. Tę wartość utrwaliły późniejsze zastosowania wojskowe i pomiarowe – prace MIT Radiation Laboratory w czasie II wojny światowej oraz specyfikacja MIL-C-17 (od 1949) – a sprzyjały jej też dostępne rozmiary rur miedzianych, z których robiono pierwsze linie (dawały ok. 51...52 Ω). Gdy liczy się wyłącznie minimum strat – jak w torach sygnału wizyjnego – standardem zostało 75 Ω. Innymi słowy: 50 Ω i 75 Ω to nie są wartości święte, ani przypadkowe, lecz dwa różne punkty na tej samej krzywej kompromisu w parametrach kabla koncentrycznego.

Ważniejszy jest wniosek pojęciowy. 50 Ω pełni rolę impedancji odniesienia systemu – wspólnej wartości, do której projektujemy generatory, kable, filtry i wejścia układów, aby łączyć je bez odbić. „Dopasowanie” znaczy: doprowadzenie do tej impedancji odniesienia, a nie „posiadanie rezystancji 50 Ω”. To rozróżnienie jest sednem kolejnego punktu.

2. Dopasowanie oznacza dopasowanie do linii, a nie „rezystancja 50 Ω”

Aby zobaczyć, czym naprawdę jest dopasowanie, trzeba myśleć falami, nie rezystancjami. Sygnał biegnący linią o impedancji charakterystycznej Z_0 napotyka na końcu obciążenie o impedancji Z . Jeśli $Z \neq Z_0$, część fali odbija się i wraca do źródła. Miarą tego odbicia jest współczynnik odbicia:

$$\Gamma = \frac{Z - Z_0}{Z + Z_0}$$

Wzór ten czyta się wprost: odbicie znika ($\Gamma = 0$) dokładnie wtedy, gdy $Z = Z_0$. Każde inne obciążenie – w tym czysto rezystancyjne 40 Ω albo 60 Ω, nie mówiąc o zespolonym – daje $\Gamma \neq 0$ i odbija część energii. To dlatego „dopasowanie” nie jest cechą samego elementu, tylko relacją między elementem a impedancją odniesienia linii.

Odbita fala nakłada się na padającą i tworzy na linii falę stojącą – stały w czasie rozkład maksimum i minimum napięcia. Jej głębokość opisuje współczynnik fali stojącej (ang. *Voltage Standing Wave Ratio*, VSWR), czyli stosunek napięcia w maksimum do napięcia w minimum:

$$VSWR = \frac{1+|\Gamma|}{1-|\Gamma|}$$

Dla idealnego dopasowania $VSWR = 1$ (brak fali stojącej). $VSWR = 2$ oznacza $|\Gamma| = 1/3$, czyli odbicie ok. 11% mocy. To samo odbicie wygodnie wyraża się też w decybelach jako tłumienność odbicia (ang. *Return Loss*, RL):

$$RL = -20 \log_{10} |\Gamma|$$

RL = 20 dB znaczy, że moc odbita jest 100 razy mniej niż od padającej ($|\Gamma| = 0,1$). Im większa tłumienność odbicia w dB, tym lepsze dopasowanie – to samo zjawisko jest opisane trzema równoważnymi liczbami: Γ , VSWR, RL.

Te trzy liczby – Γ , VSWR i RL – wygodnie odczytuje się z wykresu Smitha. Wyjaśnijmy jego ideę.

Impedancja jest liczbą zespoloną, $Z = R + jX$: ma część rzeczywistą (rezystancję R) i urojoną (reaktancję X).

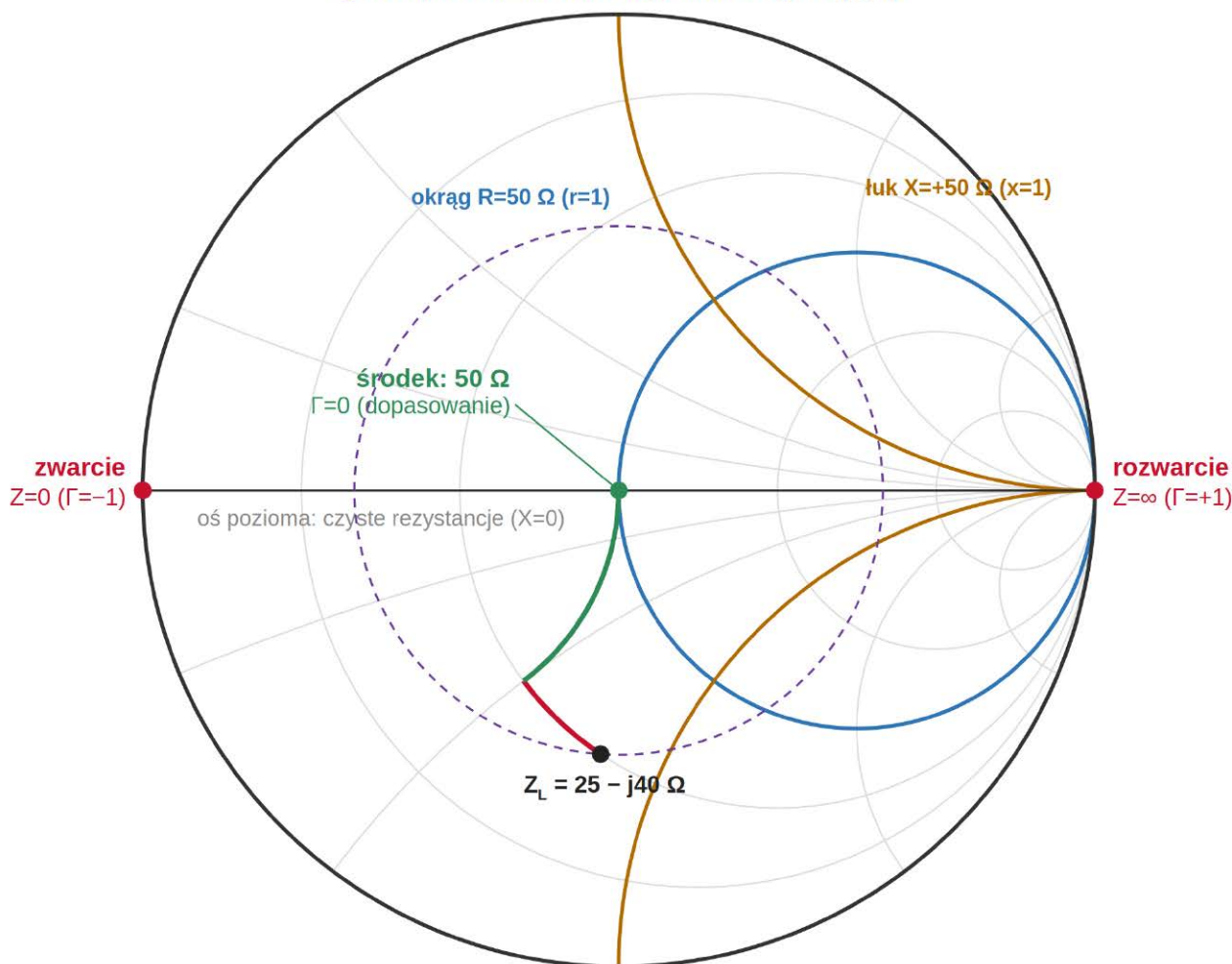
To dwie niezależne liczby o ogromnych zakresach – R od zera do nieskończoności, X od minus do plus nieskończoności. Narysowanie i „nawigowanie” po takiej nieograniczonej płaszczyźnie byłoby niewygodne.

Wykres Smitha rozwiązuje to jednym pomysłem: zamiast impedancji rysujemy współczynnik odbicia Γ . Γ też jest liczbą zespoloną, ale dla każdego obciążenia pasywnego jego moduł nie przekracza jedności ($|\Gamma| \leq 1$) – zawsze mieści się wewnątrz koła o promieniu 1. Dzięki temu cała nieskończona płaszczyzna impedancji zostaje zawarta w skończonej powierzchni koła – i to koło to właśnie wykres Smitha. Ponieważ Γ i Z wiąże jednoznaczny wzór $\Gamma = (Z - Z_0)/(Z + Z_0)$, każdej impedancji odpowiada dokładnie jeden punkt w tym kole.

Zacznijmy od punktów charakterystycznych (**rysunek 2**). Środek koła to $\Gamma=0$, czyli $Z=Z_0=50 \Omega$ – dopasowanie

Wykres Smitha: cała impedancja ściśnięta w jedno koło

górna połowa: reaktancja indukcyjna (+jX)



dolna połowa: reaktancja pojemnościowa (-jX)

— krok 1: reaktancja szeregową (po stałym R)

— krok 2: element równoległy (po stałym G)

2. Wykres Smitha odwzorowuje całą (nieskończoną) płaszczyznę impedancji na skończone koło współczynnika odbicia Γ . Środek to dopasowanie (50Ω , $\Gamma=0$), prawy skraj – rozwarcie ($Z=\infty$), lewy – zwarcie ($Z=0$); górna połowa jest indukcyjna, dolna pojemnościowa. Zaznaczono okrąg stałej rezystancji $R=50 \Omega$ ($r=1$), łuk stałej reaktancji $X=+50 \Omega$ ($x=1$) oraz dwuelementową trajektorię dopasowania $Z_L=25-j40 \Omega$ do środka

idealne; to do niego dążymy. Skrajny punkt po prawej ($\Gamma=+1$) to rozwarcie, $Z=\infty$. Skrajny punkt po lewej ($\Gamma=-1$) to zwarcie, $Z=0$. Cały zewnętrzny okrąg ($|\Gamma|=1$) to obciążenia czysto reaktancyjne ($R=0$), bezstratne – odbijają całą moc. Oś pozioma to czyste rezystancje ($X=0$); górna połowa koła odpowiada reaktancjom indukcyjnym ($+jX$), a dolna – pojemnościowym ($-jX$).

Gdyby w kole widniał tylko goły punkt Γ , nie odczytali byśmy z niego R ani X . Dlatego wykres ma naniesioną siatkę: krzywe stałej rezystancji i stałej reaktancji, przeliczone do płaszczyzny Γ . Krzywe stałego R stają się okręgami – wszystkie przechodzą przez prawy skraj (rozwarcie); im większe R , tym mniejszy okrąg bliżej prawej strony. Okrąg $R=50 \Omega$ ($r=1$) przechodzi przez środek. Krzywe stałego X to łuki wychodzące z tego samego prawego skraju, wygięte w górę (indukcyjne) lub w dół (pojemnościowe). Aby nanieść daną impedancję, znajdujemy przecięcie jej okręgu stałego R i łuku stałego X – to jest jej punkt Γ . I odwrotnie: dla dowolnego punktu odczytujemy R i X z tego, na którym okręgu i łuku leży.

Odległość punktu od środka niesie osobną, praktyczną informację: to $|\Gamma|$, czyli wielkość odbicia – przeliczalna wprost na VSWR. Punkt w środku oznacza dopasowanie idealne; im dalej od środka, tym gorsze dopasowanie. Okrąg zakreślony wokół środka łączy punkty o stałym VSWR (na rysunku 2 przerywany). Jednym rzutem oka widać więc, jak bardzo obciążenie jest niedopasowane.

Największa użyteczność wykresu ujawnia się przy projektowaniu, bo elementy obwodu przesuwają punkt po tych właśnie krzywych. Dodanie reaktancji szeregowej (cewki lub kondensatora w szereg) zmienia X , ale nie R – punkt ślizga się po swoim okręgu stałej rezystancji. Dodanie elementu równoległego najprościej odczytać na bliźniaczej siatce stałej konduktancji – wtedy punkt ślizga się po okręgu stałej konduktancji. Dlatego obwód

dwuelementowy (typu L) to dwa takie przesunięcia, które z dowolnego punktu startowego doprowadzają do środka. Dokładnie to pokazują trajektorie na rysunku 2: obciążenie $Z_L=25-j40 \Omega$ sprowadzamy do 50Ω najpierw reaktancją szeregową (ruch po okręgu stałego R), a potem elementem równoległym (ruch po okręgu stałej konduktancji). Wykres Smitha zamienia więc rozwiązywanie zespolonych równań na nakreślenie dwóch łuków do środka – dlatego jest użytecznym narzędziem, a nie ozdobą.

Trzeba jednak pamiętać o ograniczeniu: reaktancje zależą od częstotliwości, więc dane obciążenie zakreśli na wykresie krzywą, gdy zmieniamy częstotliwość, a sieć, która trafia w środek na jednej częstotliwości, schodzi z niego na innej. Dopasowanie jest więc z natury wąskopasmowe – „dopasowany” układ bywa dopasowany tylko dla warunków, dla których go zaprojektowano.

3. Dopasowanie na moc $\neq$ dopasowanie na szumy

Tu zaczyna się najpoważniejsze nieporozumienie. Intuicja podpowiada, że „dopasowanie” ma jeden cel: maksymalny transfer energii. Rzeczywiście, maksimum mocy źródło oddaje obciążeniu przy dopasowaniu sprzężonym – gdy impedancja źródła jest sprzężeniem zespolonym impedancji obciążenia:

$$Z_S = Z_L^* \Leftrightarrow \Gamma_S = \Gamma_{in}^*$$

W języku czwórników RF odpowiada to jednoczesnemu dopasowaniu sprzężonemu na wejściu i wyjściu, które daje maksymalne dostępne wzmocnienie (ang. *maximum available gain*, G_{ma}). Gdyby liczyło się tylko wzmocnienie, problem byłby zamknięty. Ale we wzmacniaczu niskosumnym (ang. *Low-Noise Amplifier*, LNA) liczy się co innego – szum. A minimum szumów wypada przy zupełnie innej impedancji źródła niż maksimum wzmocnienia.

Współczynnik szumów (ang. *Noise Figure*, NF) tranzystora nie jest stałą – zależy od tego, jaką impedancję „widzi” wejście od strony źródła. Zależność tę opisuje równanie szumów czwórnika:

$$F = F_{min} + \frac{4r_n |\Gamma_S - \Gamma_{opt}|^2}{(1 - |\Gamma_S|^2)(1 + |\Gamma_{opt}|^2)}$$

Rozłóżmy je na czynniki. F_{min} to najmniejszy osiągalny współczynnik szumów danego tranzystora. Γ_{opt} to optymalny współczynnik odbicia źródła, przy którym to minimum jest osiągalne. r_n to znormalizowana rezystancja szumów – miara tego, jak gwałtownie rośnie szum, gdy oddalamy się od Γ_{opt} . Kluczowa jest sama struktura wzoru: szum jest minimalny, gdy $\Gamma_S = \Gamma_{opt}$, i rośnie z kwadratem odległości od tego punktu. Tymczasem maksimum wzmocnienia wypada w innym punkcie – nazwijmy go Γ_{ms} . Te dwa punkty na ogół się nie pokrywają (rysunek 3).

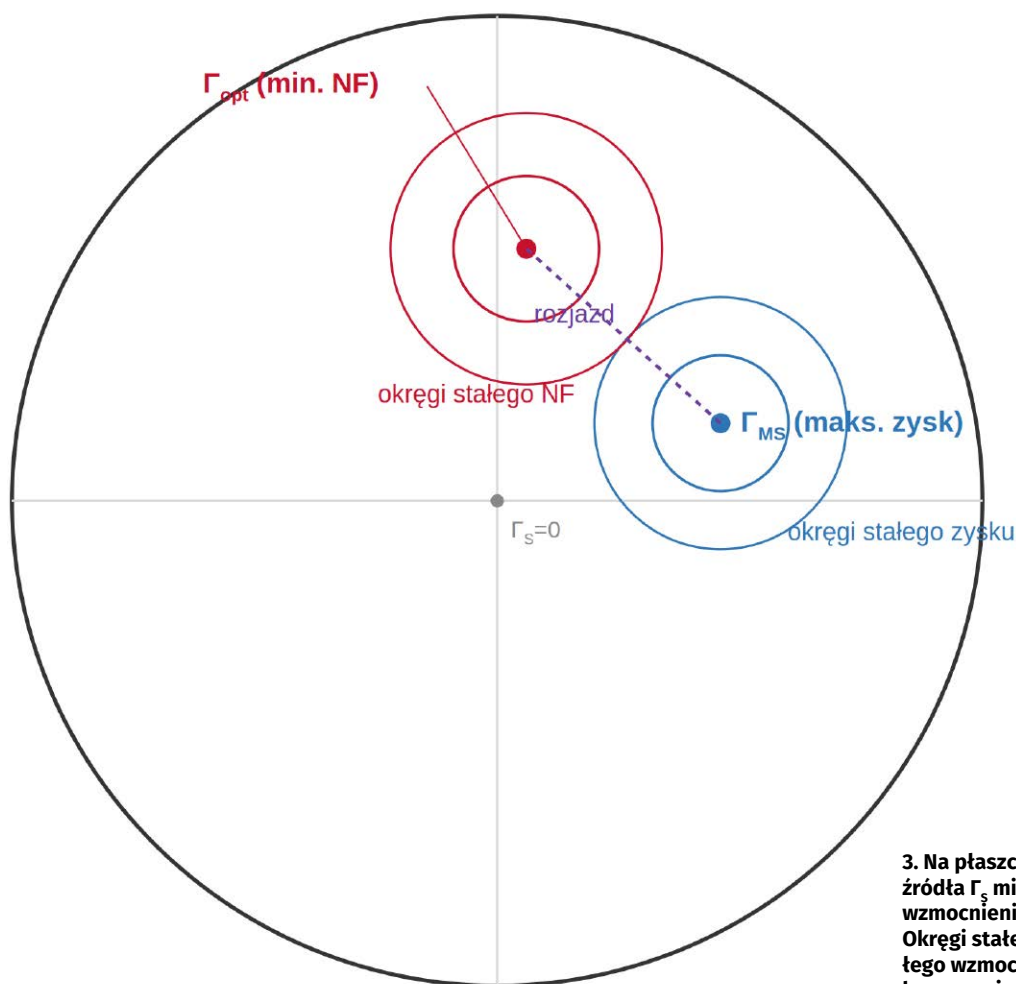
Dwa przykłady dopasowania z praktyki

Jak to wygląda przy konkretnych liczbach? Dwa typowe podejścia – skupione i rozproszone.

Przykład 1 – sieć L z elementów skupionych (poniżej ok. 1 GHz). Dopasujemy obciążenie 200Ω do toru 50Ω na 100 MHz. Dobroć sieci wynika z ilorazu rezystancji: $Q = \sqrt{(200/50 - 1)} = \sqrt{3} \approx 1,73$. Przy obciążeniu wstawiamy element równoległy, a w stronę źródła szeregowy; w wariacie dolnoprzepustowym wychodzi kondensator równoległy ok. 13,8 pF i cewka szeregową ok. 138 nH. Ta sieć dopasowuje idealnie tylko na 100 MHz – jest wąskopasmowa.

Przykład 2 – transformator ćwierćfalowy (odcinek linii). Dopasujemy antenę 100Ω do toru 50Ω na 1 GHz. Wstawiamy odcinek linii o impedancji równej średniej geometrycznej dopasowywanych wartości: $Z = \sqrt{(50 \cdot 100)} \approx 70,7 \Omega$, o długości ćwierć fali. Na mikropasku na laminacie FR-4 (efektywne $\epsilon_r \approx 3,3$) ćwierć fali na 1 GHz to ok. 41 mm. Również to dopasowanie jest wąskopasmowe – działa tam, gdzie odcinek ma dokładnie ćwierć długości fali.

Dopasowanie na moc $\neq$ na szumy



3. Na płaszczyźnie współczynnika odbicia źródła Γ_s minimum szumów (Γ_{opt}) i maksimum wzmocnienia (Γ_{MS}) leżą w różnych miejscach. Okręgi stałego NF otaczają Γ_{opt} , okręgi stałego wzmocnienia – Γ_{MS} . Projekt LNA to wybór kompromisu między nimi

Dlatego akurat w LNA poświęcamy wzmocnienie na rzecz szumów? Odpowiada za to wzór Friisa na współczynnik szumów kaskady:

$$F = F_1 + \frac{F_2 - 1}{G_1} + \frac{F_3 - 1}{G_1 G_2} + \dots$$

Szum całego toru jest zdominowany przez pierwszy stopień: wkład drugiego stopnia jest dzielony przez wzmocnienie pierwszego (G_1), trzeciego – przez iloczyn $G_1 G_2$ itd. Dlatego pierwszy wzmacniacz w torze odbiorczym projektuje się na F_{min} (dopasowanie na szumy, $\Gamma_s = \Gamma_{opt}$), nawet jeśli oznacza to nieco mniejsze wzmocnienie – bo to jego szum wyznacza poziom szumów własnych całego toru odbiorczego.

Trzy różne cele dopasowania

Fraza „Dopasowanie do 50 Ω ” nie precyzuje celu. Wzmacniacz mocy (ang. *Power Amplifier*, PA) dopasowuje się na maksimum mocy lub sprawności (metodą load-pull, nie małosygnowo). Wzmacniacz niskoszumny – na Γ_{opt} , czyli minimum szumów. Filtr – na właściwą charakterystykę w paśmie. To trzy różne „dopasowania”, które ta sama fraza zlewa w jedno.

Praktyczny wniosek: zanim powiesz „dopasowane do 50 Ω ”, powiedz – dopasowane ze względu na co: na moc, na szumy, czy na płaską charakterystykę przenoszenia.

4. „Dopasowane” $\neq$ „stabilne” (i $\neq$ „działa”)

Współczynnik odbicia opisuje jedną rzecz: ile fali wraca od portu. Nie mówi nic o tym, czy wzmacniacz jest stabilny. To rozróżnienie bywa istotne, bo wzmacniacz może mieć idealnie dopasowane wejście ($|\Gamma| \approx 0$) i mimo to oscylować.

O stabilności czwórnika decydują jego parametry rozproszenia (ang. *scattering parameters*, S) oraz impedancje, którymi go obciążymy od strony źródła i wyjścia. Ilościową miarą jest współczynnik stabilności Rolletta:

$$K = \frac{1 - |S_{11}|^2 - |S_{22}|^2 + |\Delta|^2}{2|S_{12}S_{21}|}$$

gdzie wyznacznik macierzy rozproszenia spełnia warunek dopełniający:

$$|\Delta| = |S_{11}S_{22} - S_{12}S_{21}| < 1$$

Układ jest bezwarunkowo stabilny – to znaczy nie wzbudzi się przy żadnym pasywnym obciążeniu – tylko wtedy, gdy jednocześnie $K > 1$ oraz $|\Delta| < 1$. Gdy $K < 1$, wzmacniacz jest stabilny jedynie warunkowo: istnieją pasywne impedancje źródła lub obciążenia,

„Dopasowany” ≠ stabilny: okręgi stabilności



4. Okręgi stabilności dzielą płaszczyznę Γ_s na obszar stabilny i niestabilny. Gdy $K < 1$, część impedancji „dopasowujących” leży w obszarze niestabilnym – układ dopasowany, lecz oscylujący

gdy $K < 1$, część „dopasowań” prowadzi do oscylacji — dopasowanie nie gwarantuje stabilności

przy których na wejściu lub wyjściu pojawia się rezystancja ujemna ($|\Gamma_{in}| > 1$ albo $|\Gamma_{out}| > 1$) – czyli warunek oscylacji.

Granice tych obszarów wyznaczają na wykresie Smitha okręgi stabilności (rysunek 4). Dzielą one płaszczyznę współczynnika odbicia na część stabilną i niestabilną. I tu tkwi pułapka: punkt „dopasowany do 50 Ω ” może wypaść właśnie w obszarze niestabilnym. Samo dopasowanie nie chroni przed wzbudzeniem.

Stąd praktyczne wnioski. Po pierwsze, współczynniki K i $|\Delta|$ trzeba sprawdzać w całym paśmie – i poza nim, bo oscylacje najchętniej rodzą się tam, gdzie wzmocnienie jest jeszcze duże, a nikt nie analizuje (niskie częstotliwości, poza pasmem pracy). Po drugie, jeśli układ jest tylko warunkowo stabilny, stabilizuje się go najpierw (elementy stratne, rezystory w odpowiednich miejscach), a dopiero potem dopasowuje. Kolejność odwrotna to klasyczny błąd: „dopasowałem, a piszczy”.

Wniosek dla tezy artykułu: niski $|S_{11}|$ na jednej częstotliwości nie mówi nic ani o stabilności, ani o zachowaniu układu poza tą częstotliwością. „Dopasowany” to nie to samo co „działający prawidłowo”.

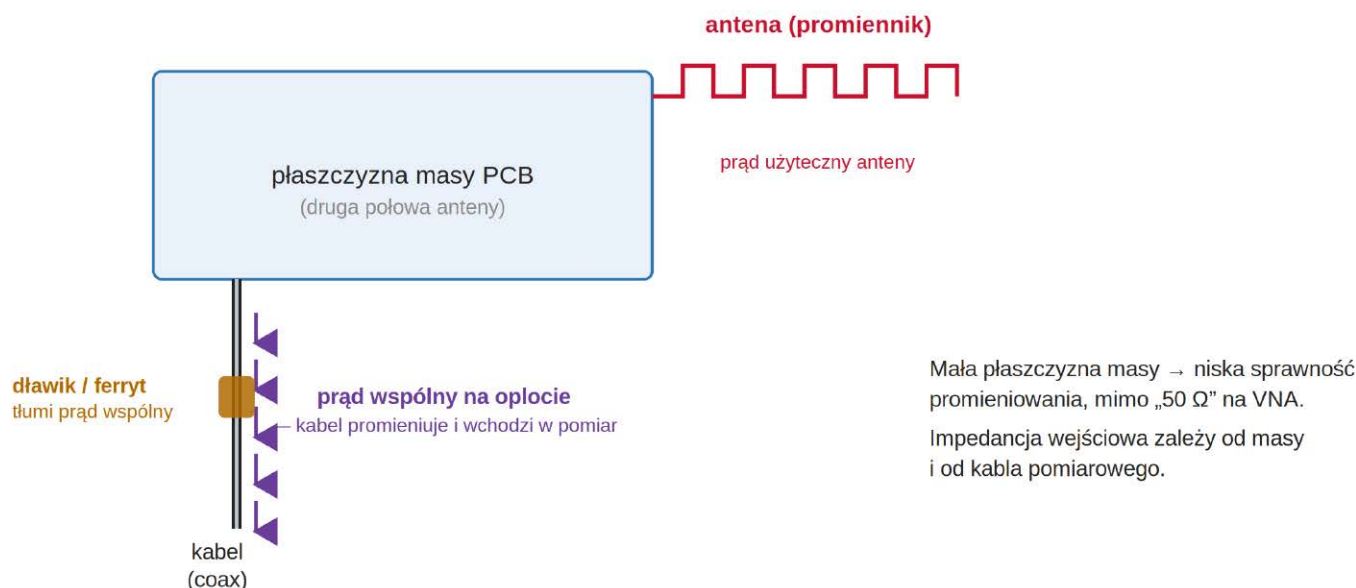
5. Antena na płycie: płaszczyzna masy jest częścią anteny

Przenieśmy się do anten drukowanych – bo tu „dopasowanie do 50 Ω ” myli najdotkliwiej. Weźmy zadanie z życia: małą antenę wtopioną w płytkę urządzenia bezprzewodowego, na przykład modułu Bluetooth, pilota albo czujnika IoT. To antena użytkowa – chcemy, żeby promieniowała. Dostrajamy ją, na analizatorze widzimy piękne dopasowanie do 50 Ω , a mimo to zasięg jest fatalny. Dlaczego „dopasowana” antena może nie promieniować?

Zacznijmy od skali, bo w niej tkwi sedno. Fala o częstotliwości 2,4 GHz ma długość około 12,5 cm. Antena na płycie miewa 1 cm – jest wielokrotnie mniejsza od fali, którą ma wypromieniować. Takie anteny (w żargonie: elektrycznie małe, czyli dużo mniejsze od długości fali) są z natury mało sprawne i zachowują się przewrotnie.

Najważniejsze jest to, że taka antena nie działa sama. Potrzebuje drugiej połowy – powierzchni, względem której „odpycha” prąd. Rolę tej drugiej połowy pełni miedziana płaszczyzna masy płytki (w żargonie nazywana przeciwwagą). Prąd wielkiej częstotliwości nie płynie bowiem tylko w samym elemencie anteny: wypływa

Mała antena na płycie: masa i kabel też promieniują



5. Mała antena na płycie potrzebuje drugiej połowy – miedzianej płaszczyzny masy, przez którą wraca prąd i która też promieniuje. Gdy masa jest za mała, prąd wchodzi na zewnętrzną stronę ekranu kabla (prąd wspólny): kabel promieniuje i zaburza pomiar. Dlatego dobre dopasowanie do 50 Ω nie gwarantuje dobrego promieniowania; pomaga większa płaszczyzna masy i dławik na kablu

z niego i wraca przez płaszczyznę masy – a skoro płynie w masie, to masa też promieniuje. Płaszczyzna masy nie jest więc biernym podłożem; jest drugim biegunem anteny i promieniuje na równi z samym elementem. Stąd prosty wniosek: jeśli płaszczyzna masy jest za mała, antena ma za małą drugą połowę i cały układ promieniuje słabo – nawet gdy analizator pokazuje idealne 50 Ω.

I tu jest sedno tego tematu. Dopasowanie mówi wyłącznie jedno: że moc wchodzi do zacisków anteny, nie odbijając się z powrotem. Nie mówi, co się z tą mocą dalej dzieje. A może ona pójść w dwie strony: albo zostać wypromieniona jako fala (o to nam chodzi), albo zamienić się w ciepło w rezystancji strat (czysta strata). Analizator tego nie odróżnia – dla niego oba przypadki wyglądają tak samo: „moc weszła, nic się nie odbiło”. Dowód skrajny: rezystor 50 Ω jest dopasowany idealnie i nie promieniuje nic, bo cała moc grzeje rezystor. To, jak dobrym źródłem fal jest antena, opisuje sprawność promieniowania (stosunek mocy wypromienionej do dostarczonej), a nie tłumienność odbicia. Wykres $|S_{11}|$ tej różnicy po prostu nie widzi.

Do tego dochodzi druga, osobna komplikacja – tym razem po stronie pomiaru. Gdy podłączamy antenę do analizatora kablem współosiowym, a płaszczyzna masy jest za mała, antena „szuka” brakującej drugiej połowy i znajduje ją w kablu. Prąd zaczyna płynąć po zewnętrznej stronie ekranu (oplotu) kabla. Taki prąd – płynący wspólnie po powierzchni ekranu, dlatego nazywany prądem wspólnym (ang. *common-mode*) – sprawia, że kabel staje się przedłużeniem anteny i sam promieniuje.

Psuje to sprawę na dwa sposoby. Po pierwsze, fałszuje pomiar: mierzona impedancja wejściowa zaczyna zależeć

od długości i ułożenia kabla. Jeśli przesunięcie kabla zmienia odczyt – to nie antena się zmieniła, tylko mierzymy antenę razem z kablem. Po drugie, w gotowym urządzeniu ten sam prąd wspólny oznacza niekontrolowane promieniowanie z kabli i ścieżek – typowe źródło problemów z kompatybilnością elektromagnetyczną.

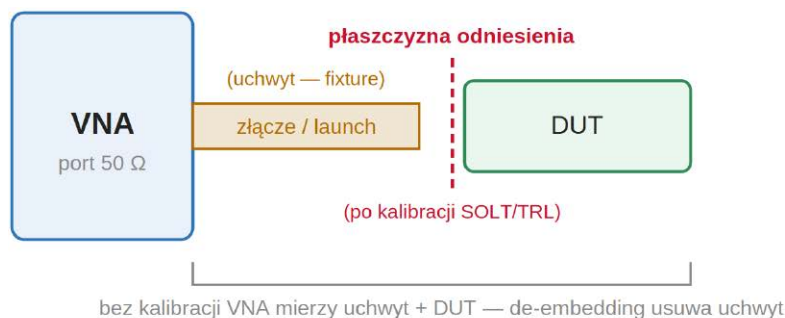
Lekarstwo jest konkretne. Zapewnić odpowiednio dużą płaszczyznę masy (dużą drugą połowę anteny). Założyć na kabel, blisko anteny, dławik prądu wspólnego – koralek ferrytowy lub balun – który blokuje prąd na zewnętrznej stronie ekranu. I mierzyć antenę w docelowym otoczeniu: z tą masą, obudową i zasilaniem, w których będzie pracować. Wniosek zgodny z tezą całego artykułu: „dostrojona do 50 Ω” nie znaczy „dobrze promieniuje”. Sprawność promieniowania to osobny parametr, którego dopasowanie nie gwarantuje.

6. VNA „widzi” co innego, niż myślisz

Skoro pomiar bywa mylący, przyjrzyjmy się narzędziu, któremu ufamy najbardziej – wektorowemu analizatorowi sieci (ang. *Vector Network Analyzer*, VNA). VNA nie mierzy „prawdy absolutnej”, lecz wartości względem dwóch założeń: płaszczyzny odniesienia i impedancji odniesienia. Oba potrafią wprowadzić w błąd (rysunek 6).

Pierwsza pułapka to płaszczyzna odniesienia i uchwyt pomiarowy. Surowy pomiar obejmuje złącza, przejścia i uchwyt (ang. *fixture*), a nie tylko badany element (ang. *Device Under Test*, DUT). Kalibracja – metodą SOLT (zwarcie–rozwarcie–obciążenie–połączenie) lub TRL (linia–odbicie–linia) – przesuwa płaszczyznę odniesienia do zacisków DUT. Bez poprawnej kalibracji i de-embeddingu

VNA „widzi” co innego: płaszczyzna odniesienia i renormalizacja



Renormalizacja:

ten sam DUT ma inne S-parametry przy różnej impedancji odniesienia.

pomiar w 50 Ω: S_{11} „dobre”
ale filtr pracuje między stopniami 200 Ω

tłumik 50 Ω: S_{11} idealne, zysk zerowy
niski $|S_{11}| \neq$ układ działa

6. VNA mierzy względem płaszczyzny odniesienia (ustalanej kalibracją) i impedancji odniesienia (zwykle 50 Ω). Bez de-embeddingu pomiar obejmuje uchwyt; po renormalizacji ten sam element ma inne S-parametry. Niski $|S_{11}|$ tłumika dowodzi, że dopasowanie nie oznacza użyteczności

mierzysz uchwyt razem z elementem i przypisujesz elementowi cudze właściwości.

Druga pułapka to renormalizacja. S-parametry są zdefiniowane względem impedancji odniesienia – domyślnie 50 Ω. Ten sam DUT ma inne S-parametry, gdy odniesimy je do innej impedancji. Filtr o świetnym „50-omowym” S_{11} może zachowywać się zupełnie inaczej, gdy w rzeczywistym torze pracuje między stopniami o impedancji np. 200 Ω. Aby przewidzieć jego pracę, trzeba renormalizować S-parametry do rzeczywistych impedancji układu (albo od razu projektować i mierzyć przy nich). Odbicie na styku dwóch światów impedancji opisuje ten sam wzór, co na początku artykułu:

$$\Gamma_{int} = \frac{Z_{sys} - Z_0}{Z_{sys} + Z_0}$$

Trzecia, najkrótsza lekcja: niski $|S_{11}|$ jest warunkiem koniecznym, ale nie wystarczającym. Tłumik 50 Ω ma niemal idealne S_{11} i – poza tłumieniem – nie robi nic. Dobre dopasowanie oznacza „mało odbić”, a nie „układ realizuje swoje zadanie”. Do oceny działania służą inne parametry: S_{21} (transmisja), współczynnik szumów, sprawność promieniowania, moc, liniowość.

I domknięcie wątku z poprzedniego rozdziału: przy pomiarze małych anten prąd wspólny na kablu pomiarowym zmienia odczyt, gdy poruszamy kablem. To sygnał, że mierzymy tor razem z kablem, a nie samą antenę – stąd dławik na kablu pomiarowym bywa warunkiem sensownego pomiaru, nie ozdobą. Te subtelności kalibracji, de-embeddingu i renormalizacji to codzienność metrologii mikrofalowej opisanej m.in. przez Joela Dunsmore’a.

Nie tylko 50 Ω

50 Ω to jeden z wielu standardów impedancji – wybierany zależnie od celu, nie „jedynie słuszny”:

50 Ω – RF, radary, aparatura pomiarowa; standaryzacja utrwalona wojskową specyfikacją MIL-C-17 (1949) i pracami MIT Radiation Laboratory z czasów wojny.

75 Ω – wideo, broadcast i długie tory koncentryczne – wybór minimum strat.

93...100 Ω – dawna instrumentacja (93 Ω) oraz tory różnicowe 100 Ω (pary skręcane, Ethernet).

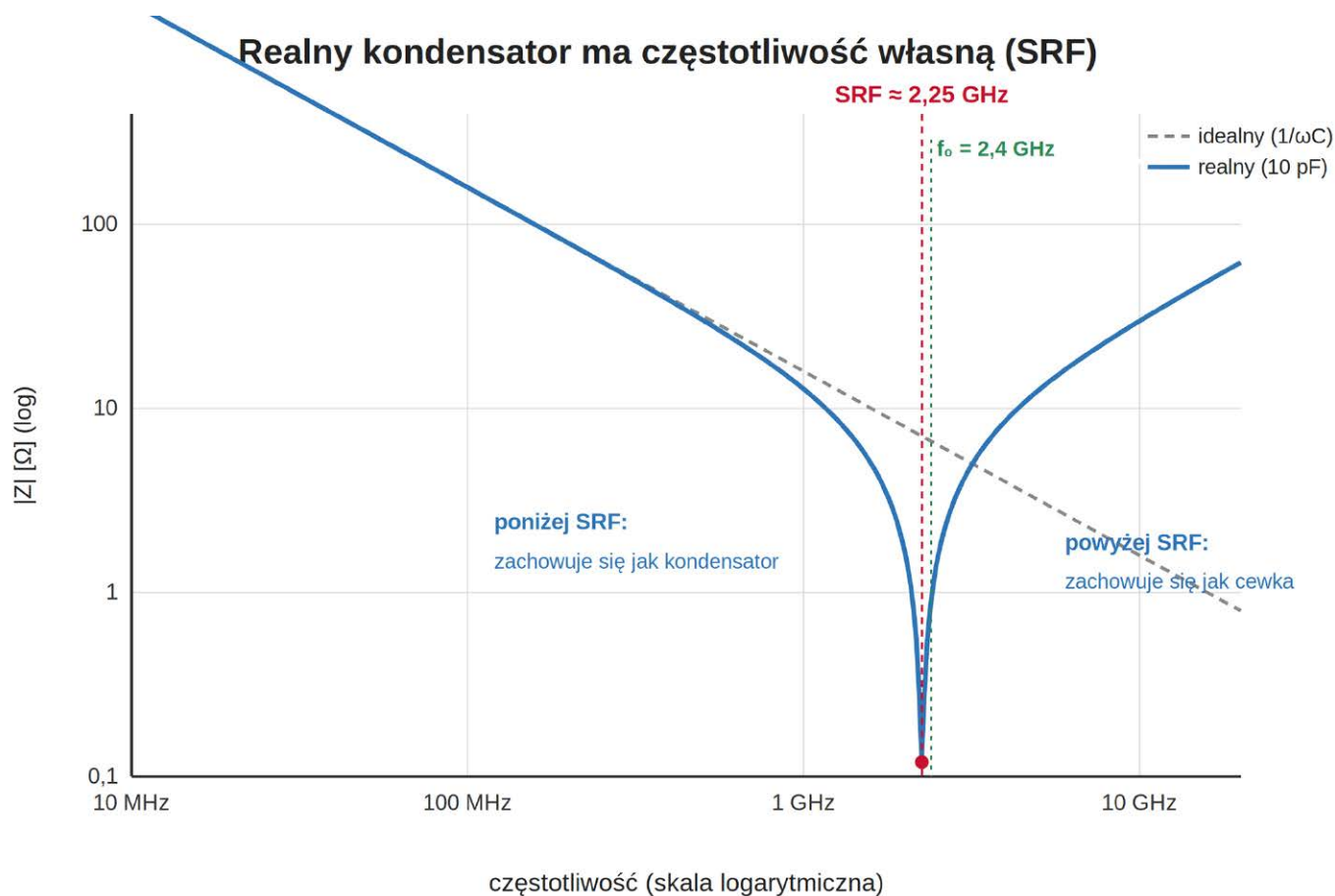
300...600 Ω – linie symetryczne (twin-lead 300 Ω, ladder-line) i historyczne 600 Ω w telefonii oraz starym torze audio.

Uwaga pojęciowa: w RF dopasowujemy impedancje (transfer mocy, brak odbić). W torze audio małosygnałowym celowo dajemy impedancję wejściową znacznie większą od wyjściowej (mostkowanie napięciowe), by nie obciążać źródła – „dopasowanie na moc” byłoby tam błędem. To samo słowo, inny cel.

7. Dlaczego skopiowany layout referencyjny 1:1 przestaje działać

Częsta sytuacja: bierzemy gotowy layout referencyjny – płytkę ewaluacyjną albo notę aplikacyjną producenta – przepisujemy wartości elementów co do sztuki i... nie działa. Dopasowanie wychodzi w innym miejscu, filtr ma inne pasmo, wzmacniacz się wzbudza. Skąd ta różnica, skoro wartości są te same? Odpowiedź ma dwie części: elementy i geometria.

Zacznijmy od elementów, bo przy wielkiej częstotliwości element nie jest tym, co jest napisane na jego obudowie. Każdy rzeczywisty kondensator ma pasożytniczą indukcyjność szeregową (wyprowadzenia, okładki), więc naprawdę jest szeregowym połączeniem pojemności C, małej indukcyjności L_s i rezystancji strat ESR. Jego impedancja wynosi:



7. Realny kondensator 10 pF ma częstotliwość własną (SRF): poniżej niej działa jak kondensator ($|Z|$ maleje), przy SRF impedancja spada do ESR, a powyżej rośnie jak u cewki. Element o innym SRF przy tej samej częstotliwości pracy (f_0) może być już po drugiej stronie rezonansu

$$Z(\omega) = ESR + j(\omega L_s - \frac{1}{\omega C})$$

Poniżej pewnej częstotliwości dominuje człon $1/\omega C$ – element jest kondensatorem, a jego impedancja maleje z częstotliwością. Przy częstotliwości własnej (ang. *self-resonant frequency*, SRF) reaktancja indukcyjna i pojemnościowa znoszą się, a impedancja spada do samego ESR:

$$f_{SRF} = \frac{1}{2\pi\sqrt{L_s C}}$$

Powyżej SRF dominuje już ωL_s : kondensator zachowuje się jak cewka (rysunek 7). Dokładnie to samo dotyczy cewek – mają pasywną pojemność i własną SRF, powyżej której stają się pojemnościowe.

Konsekwencja dla „kopii 1:1” jest wprost: ten sam nominal 10 pF w innej obudowie albo od innego producenta ma inne L_s , a więc inne SRF. Przy Twojej częstotliwości pracy może już leżeć po drugiej stronie rezonansu – zachowywać się indukcyjnie zamiast pojemnościowo. Wartość się zgadza, zachowanie nie.

Druga, równie ważna przy czynna to geometria. Przy wielkiej częstotliwości ścieżka na płytce nie jest

zwykłym połączeniem – krótki jej odcinek to element o własnej indukcyjności i pojemności, czyli kawałek linii transmisyjnej. Jego impedancja i wnoszone przesunięcie fazy zależą od długości, szerokości ścieżki, stałej dielektrycznej laminatu (E_r) i wysokości nad płaszczyzną masy (układ warstw, ang. *stack-up*). W layoucie referencyjnym te ścieżki są częścią sieci dopasowania – projektant dostrzegł je razem z elementami. Jeśli przepiszysz wartości, ale poprowadzisz ścieżki inaczej (dłużej, wężej, na innym

REKLAMA

Hurtownia elementów elektronicznych „AKSOTRONIK” zaprasza do swojego sklepu internetowego
Zaloguj się i kupuj ON-LINE na naszej stronie:
WWW.AKSOTRONIK.COM.PL

Magnezy neodymowe
oraz ferrytowe
Ceny od 0.10zł

Przełączniki klawiszowe
wodoszczelne/pyłoszczelne
Ceny od 2.40zł

Druty oporowe
od 0.16 do 0.81mm
Ceny od 5.70zł

Prowadniki do przewodów
Ceny od 11.00zł

Kostki elektryczne
zaciskowe
Ceny od 0.22zł

Szczotki węglowe
do elektronarzędzi
Ceny od 2.60zł/kpl

Przełączniki do elektronarzędzi
zwykłe i elektromagnetyczne
Ceny od 7.00zł

Złącza hermetyczne
Supersal
Ceny od 1.10zł/kpl

Pudełka/organizery
Ceny od 0.95zł

Zestawy śrubek M2, M3
z nakrętkami i podkładkami
Ceny od 2.50zł

Aksotronik
ELEMENTY ELEKTRONICZNE

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

laminacie), zmieniłeś sieć dopasowania – mimo że schemat wygląda tak samo.

Dochodzi jeszcze droga powrotu prądu i indukcyjność przelotek. Prąd wielkiej częstotliwości wraca tuż pod ścieżką, w płaszczyźnie masy; przerwa w masie albo przeniesiona przelotka zmienia pętlę prądu i dokładnie indukcyjność. To też jest część sieci – niewidoczna na schemacie.

Praktyczny wniosek: layout referencyjny da się powtórzyć tylko wtedy, gdy skopiujesz całość – geometrię ścieżek, układ warstw i elementy o tej samej SRF (ta sama obudowa, ten sam producent) – a nie same wartości. Jeśli którykolwiek z tych warunków się zmienia, dopasowanie trzeba przestroić na własnej płytce, mierząc analizatorem. „Te same wartości” to nie „ten sam układ”.

8. W pracowni

Praktyczne uwagi dotyczące kolejności pracy, które oszczędzą wielu nieporozumień z „50 omami”. Przejdźmy ją krok po kroku, wyjaśniając każde pojęcie.

Krok 1: nazwij cel, zanim cokolwiek dopasujesz. Od celu zależy, do czego w ogóle dopasowujesz. Jeśli chcesz maksymalnej mocy sygnału, celujesz w dopasowanie sprzężone (impedancja źródła równa sprzężeniu zespolonemu obciążenia). Jeśli budujesz wzmacniacz niskoszumny, celujesz w Γ_{opt} tranzystora – punkt leżący gdzieś indziej niż maksimum wzmocnienia (pkt. 3). We wzmacniaczu mocy najlepszemu obciążeniu nie wyliczysz jednym wzorem – znajduje się je metodą load-pull.

Czym jest load-pull? To pomiar (albo symulacja), w którym celowo zmienia się impedancję obciążenia „widzianą” przez tranzystor mocy i dla każdego ustawienia mierzy się moc wyjściową oraz sprawność. Powstaje z tego mapa krzywych stałej mocy na wykresie Smitha; z niej odczytuje się obciążenie dające największą moc – zwykle inne niż dopasowanie sprzężone.

Krok 2: ustal płaszczyznę odniesienia. To dokładnie ten punkt, który ogłaszasz „wejściem” swojego układu – miejsce, w którym pomiar ma się „liczyć”. Wszystko, co leży między przyrządem a tą płaszczyzną – kable, złącza, uchwyt – trzeba później usunąć z wyniku, inaczej zmierzysz co innego niż myślisz.

Krok 3: skalibruj analizator. Przyrząd mierzy na swoim złączu, ale kable i przejścia dokładają tłumienie oraz przesunięcie fazy. Kalibracja polega na dołączeniu wzorców o znanych właściwościach, z których analizator wylicza swoje błędy i odejmuje je w obliczeniach, przenosząc pomiar do płaszczyzny odniesienia. Metoda SOLT używa czterech wzorców: zwarcia (ang. *Short*), rozwarcia (Open), obciążenia dopasowanego (Load) i połączenia prostego (Thru). Metoda TRL (ang. *Thru-Reflect-Line*) korzysta z połączenia prostego, elementu silnie odbijającego i odcinka linii; bywa wygodniejsza na płytce, gdzie

trudno o dokładne wzorce SOLT. Po kalibracji analizator podaje impedancję widzianą w płaszczyźnie odniesienia, a nie na swoim panelu.

Krok 4: zde-embeduj uchwyt (ang. *de-embedding*). Nawet po kalibracji do złącza zostaje zwykle kawałek uchwyty pomiarowego – odcinek ścieżki, gniazdo, przejście – między skalibrowaną płaszczyzną a samym badanym elementem. De-embedding to obliczeniowe „odjęcie” znanego wpływu tego uchwyty: znając jego parametry (z modelu albo z osobnego pomiaru), analizator usuwa jego udział z wyniku, tak by został sam element. Innymi słowy – zmierzyles „element plus uchwyt”, a de-embedding zostawia sam „element”. Operacja odwrotna, czyli doliczenie wpływu uchwyty (na przykład w symulacji), nazywa się embedding.

Krok 5: renormalizuj parametry S, jeśli trzeba. Parametry rozproszenia (S) są zawsze podane względem jakiejś impedancji odniesienia – domyślnie 50 Ω . Jeśli Twój układ pracuje przy innej impedancji, na przykład 75 Ω w torze wizyjnym albo 100 Ω różnicowo w układzie scalonym, parametry podane dla 50 Ω nie opisują wprost jego zachowania. Renormalizacja to przeliczenie parametrów S tak, jakby impedancją odniesienia była rzeczywista impedancja układu – ten sam element, ale inna „miarka”.

Krok 6: sprawdź stabilność, zanim dopasujesz. Policz współczynnik Rolletta K oraz wyznacznik $|\Delta|$ w całym

Co zapamiętać

- 50 Ω to historyczny kompromis (minimum strat ~77 Ω , maksimum mocy ~30 Ω), a nie stała fizyczna, ani uniwersalne optimum.
- „Dopasowanie” znaczy brak odbicia na linii o danej impedancji ($\Gamma \rightarrow 0$), a nie „posiadanie rezystancji 50 Ω ”.
- Dopasowanie na moc (Γ_{MS}) to nie dopasowanie na szumy (Γ_{opt}); szum całego toru ustala pierwszy stopień (Friis).
- Niski $|S_{11}|$ jest konieczny, lecz niewystarczający: nie mówi o stabilności, promieniowaniu ani o zachowaniu poza częstotliwością pomiaru. „Dopasowany” $\neq$ „działa”.

Checklista projektanta

- ☒ Jaki jest cel: moc, szum, napięcie, stabilność, EMC? Dobierz do niego rodzaj dopasowania.
- ☒ Jaka jest impedancja odniesienia toru i aparatury? Czy 50 Ω wynika z interfejsu, czy z przyzwyczajenia?
- ☒ Wąsko- czy szerokopasmowo? Jaki VSWR/oddbicie jest jeszcze akceptowalny dla całego systemu?
- ☒ Czy analizator jest skalibrowany, uchwyt zde-embeddowany, a S-parametry renormalizowane do impedancji układu?
- ☒ Czy $K > 1$ i $|\Delta| < 1$ w paśmie i poza nim? Jeśli nie – najpierw stabilizacja.
- ☒ Czy elementy mają właściwe SRF przy częstotliwości pracy i czy geometria ścieżek odpowiada referencji?
- ☒ Antena: dość duża masa, dławik na kablu, pomiar w obudowie; oceniaj sprawność promieniowania, nie tylko $|S_{11}|$.

paśmie i poza nim. Warunek $K > 1$ przy $|\Delta| < 1$ oznacza stabilność bezwarunkową – żadne pasywne źródło ani obciążenie nie wywoła drgań. Jeśli $K < 1$, część „dopasowań” doprowadzi układ do wzbudzenia; najpierw go ustabilizuj (rezystory szeregowo lub równoległe, sprzężenie), a dopiero potem dopasuj. Sprawdzaj także poza pasmem pracy – tranzystor ma największe wzmocnienie przy niskich częstotliwościach i to właśnie tam najchętniej oscyluje.

Krok 7: dopasuj świadomie na wykresie Smitha. Nanieś zmierzone obciążenie i sprowadź je do środka (50Ω) sieci dopasowującą – dwuelementową typu L albo trójelementową typu Π lub T (więcej elementów daje większą kontrolę nad pasmem). Pamiętaj o dwóch pułapkach z poprzednich rozdziałów: dopasowanie jest wąskopasmowe (pkt. 2), a elementy mają swój SRF (pkt. 7). Dlatego dostrajaj na docelowej płytce, mierząc, a nie „z noty aplikacyjnej”.

Krok 8: przy antenach myśl o promieniowaniu, nie o samym dopasowaniu. Zapewnij dostatecznie dużą płaszczyzną masy, załóż dławik prądu wspólnego na kablu i mierz w docelowej obudowie. Oceniaj sprawność promieniowania, a nie sam współczynnik odbicia (pkt. 5).

Krok 9: nie uważaj, że niska wartość S_{11} gwarantuje, że układ działa. Niski współczynnik odbicia jest

Eksperci, których opinie uwzględniono w artykule

Chris Bowick – amerykański inżynier RF, autor klasycznego podręcznika *RF Circuit Design* (wyd. 1: 1982; wyd. 2: 2008), jednej z najczęściej używanych praktycznych pozycji o sieciach dopasowania, filtrach i wzmacniaczach wielkiej częstotliwości.

Joel P. Dunsmore – wieloletni inżynier i badacz związany z firmą Keysight Technologies (dawniej Agilent/Hewlett-Packard), specjalista metrologii mikrofalowej i pomiarów wektorowym analizatorem sieci; autor monografii *Handbook of Microwave Component Measurements* i współtwórca wielu technik kalibracji i de-embeddingu VNA.

Guillermo Gonzalez – profesor inżynierii elektrycznej (University of Miami), autor podręcznika *Microwave Transistor Amplifiers: Analysis and Design* (1984, 1997) – standardowej pozycji o projektowaniu wzmacniaczy metodą parametrów rozproszenia, o stabilności (współczynnik Rolletta) i dopasowaniu szumowym.

warunkiem koniecznym, ale nie wystarczającym. O działaniu mówią inne parametry: transmisja S_{21} (wzmocnienie albo tłumienie), współczynnik szumów, sprawność promieniowania, moc i liniowość, a nade wszystko stabilność. Dopasowanie to jeden z warunków poprawnego działania – nie cały projekt.

WM, DS

Bibliografia

- [1] C. Bowick, *RF Circuit Design*, wyd. 2. Newnes/Elsevier, 2008. ISBN 978-0-7506-8518-4.
- [2] G. Gonzalez, *Microwave Transistor Amplifiers: Analysis and Design*, wyd. 2. Prentice Hall, 1997. ISBN 978-0-13-254335-4.
- [3] D. M. Pozar, *Microwave Engineering*, wyd. 4. Wiley, 2011. ISBN 978-0-470-63155-3.
- [4] J. P. Dunsmore, *Handbook of Microwave Component Measurements: with Advanced VNA Techniques*, wyd. 2. Wiley, 2020. ISBN 978-1-119-47713-6.
- [5] C. A. Balanis, *Antenna Theory: Analysis and Design*, wyd. 4. Wiley, 2016. ISBN 978-1-118-64206-1.
- [6] J. M. Rollett, „Stability and Power-Gain Invariants of Linear Twoports”, *IRE Trans. Circuit Theory*, t. 9, nr 1, s. 29–32, 1962. DOI 10.1109/TCT.1962.1086854.
- [7] H. T. Friis, „Noise Figures of Radio Receivers”, *Proc. IRE*, t. 32, nr 7, s. 419–422, 1944. DOI 10.1109/JRPROC.1944.232049.
- [8] MIL-C-17 – specyfikacja kabli współosiowych (impedancje 50Ω i 75Ω). Departament Obrony USA, od 1949 (geneza standaryzacji 50Ω , wraz z pracami MIT Radiation Laboratory).
- [9] L. Espenschied, H. A. Affel, „Concentric conducting system”, patent US 1 835 031, 1931 (kabel współosiowy, Bell Telephone Laboratories; zgłoszenie 1929).

REKLAMA



przejrzyj i kupisz na stronie
www.ulubionykiosk.pl

PROJEKTY dla elektroników

- ▶ Pico Gamer – konsola do gier w stylu retro
- ▶ Dzielnik częstotliwości wzorcowej 10 MHz
- ▶ Cyfrowy terminal wideo z Raspberry Pi Pico, część 1

DIY dla wszystkich

- ▶ Termometr cyfrowy z modułem Arduino Uno
- ▶ System monitorowania warunków uprawy pszenicy
- ▶ Zaawansowany wykrywacz alkoholu w wydychanym powietrzu
- ▶ Edukacja w EdW dla szkół i uczelni: Wykład 44: Płyta CD audio

TUTORIALE

- ▶ Nauka elektroniki z modułem ESP32, część 1. Pierwsze kroki
- ▶ Kick Start część 11: pomiary parametrów środowiska. Wprowadzenie do czujnika BME280
- ▶ Audio OUT: Trzaski powodowane przez potencjometry, część 2
- ▶ Chirurgia obwodowa: Rezystancja sterowana elektronicznie, część 1



Stan naładowania to nie napięcie

Dlaczego woltomierz nie powie Ci, ile energii zostało w ogniwie – i co od dwudziestu lat robią z tym problemem twórcy algorytmów zarządzania baterią

Odruch jest niemal automatyczny: bierzemy woltomierz, przykładamy go do ogniwa i z odczytu „wnioskujemy”, ile w nim zostało. Na pewnym poziomie to działa – rozładowane ogniwo ma niższe napięcie niż naładowane. Problem w tym, że związek między napięciem a stanem naładowania jest słaby, niejednoznaczny i – dla najpopularniejszej dziś chemii litowo-żelazowo-fosforanowej (ang. Lithium Iron Phosphate, LFP) – przez większość zakresu pracy wręcz nieczytelny. Stąd teza tego artykułu: napięcie nie jest wskaźnikiem stanu naładowania (ang. State of Charge, SoC). Co więcej, SoC nie jest nawet wielkością mierzalną bezpośrednio – jest stanem wewnętrznym ogniwa, który trzeba estymować. Wokół tej jednej obserwacji wyrosła cała inżynieria nowoczesnych systemów BMS (ang. Battery Management System).

1. Czym właściwie jest SoC (a czym nie jest)

Stan naładowania definiujemy jako stosunek ładunku aktualnie dostępnego do ładunku, jaki ogniwo może pomieścić między umownymi granicami pełnego i pustego. Formalnie jest to wielkość całkowita: bieżący SoC to wartość początkowa pomniejszona o scałkowany w czasie prąd, znormalizowany przez pojemność ogniwa.

$$z(t) = z(t_0) - \frac{1}{Q} \int_{t_0}^t \eta i(\tau) d\tau$$

gdzie z oznacza SoC, Q – pojemność ogniwa, η – sprawność kulombową, a $i(\tau)$ – prąd (dodatni przy rozładowaniu). Już sama postać tej definicji jest pouczająca: po prawej stronie nie ma napięcia. SoC jest całką z prądu, a nie odczytem woltomierza.

To prowadzi do punktu, który Gregory L. Plett – profesor University of Colorado Colorado Springs (UCCS), autor dwutomowej monografii o systemach BMS i twórca metody estymacji SoC opartej na rozszerzonym filtrze Kalmana – powtarza w swoich pracach i wykładach: stan naładowania nie jest wielkością, którą da się zmierzyć czujnikiem. Można go wyłącznie oszacować – na podstawie tego, co da się zmierzyć (napięcia, prądu,

temperatury) oraz modelu, który te wielkości wiąże ze stanem wewnętrznym [1, 2].

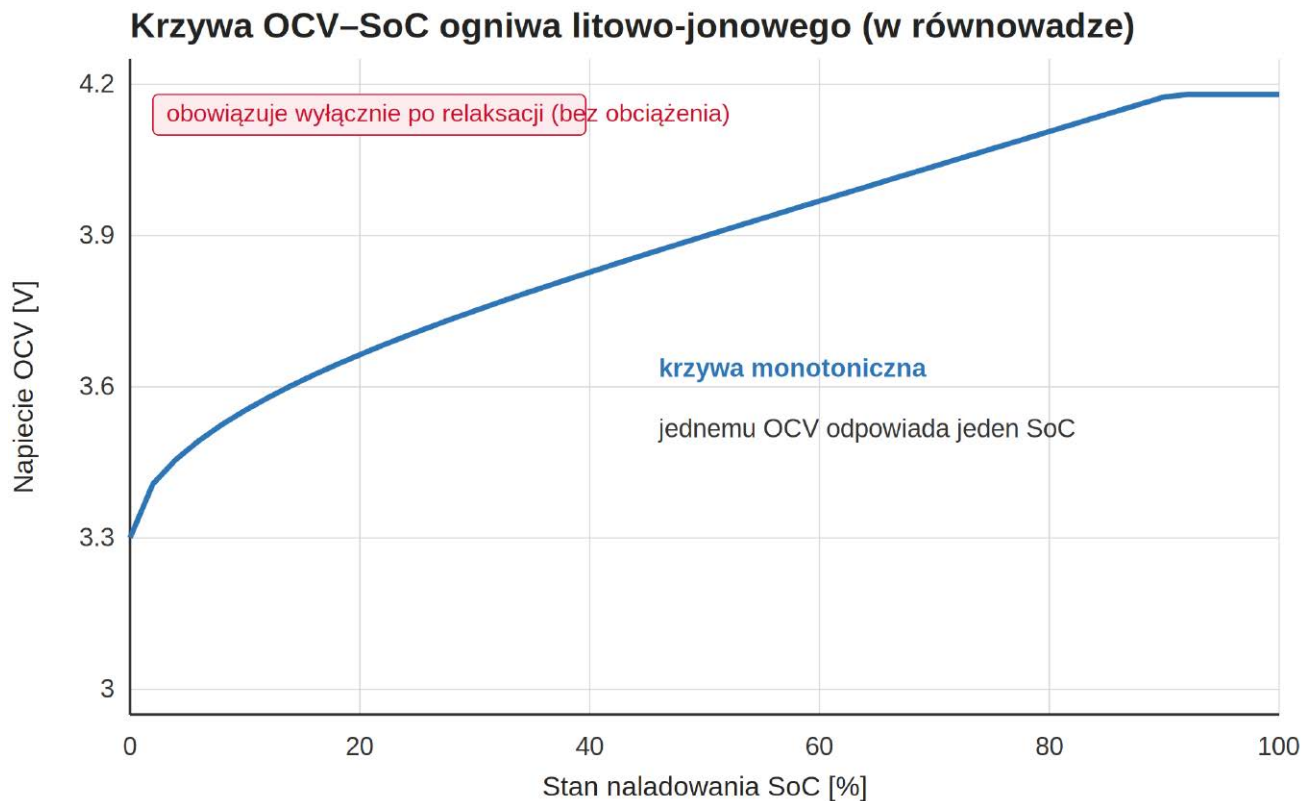
Skąd zatem rozpowszechnione „tabelki napięcie–SoC”? Biorą się z faktu, że w stanie równowagi – to znaczy po odpowiednio długim odpoczynku ogniwa, bez prądu – istnieje monotoniczna zależność między napięciem rozwartego obwodu (ang. Open-Circuit Voltage, OCV) a stanem naładowania (rysunek 1). W tym i tylko w tym sensie napięcie „niesie” informację o SoC. Ujednolicony model OCV dla ogniw litowo-jonowych opisali m.in. Weng, Sun i Peng [3].

Kluczowe rozróżnienie brzmi więc tak: OCV to nie to samo co napięcie zaciskowe. OCV jest napięciem ogniwa

Okiem eksperta – Gregory L. Plett (UCCS)

W wywiadzie dla Artech House (2024) Plett i M. Scott Trimboli zwracają uwagę, że stosowane dziś w BMS modele zastępcze (ang. Equivalent-Circuit Model, ECM) bardzo dobrze przewidują napięcie ogniwa – i dlatego oparte na nich metody potrafią dobrze estymować SoC oraz stan energii. Ich ograniczeniem jest jednak to, że jako modele empiryczne nie odtwarzają wewnętrznych zmiennych elektrochemicznych ogniwa, a więc nie mogą przewidzieć warunków prowadzących do przyspieszonej degradacji [4].

Wniosek dla projektanta: dobra estymacja SoC i dobra estymacja „zdrowia” ogniwa to dwa różne zadania, opierające się na różnych modelach.



1. Zależność OCV–SoC ogniwa litowo-jonowego w stanie równowagi. Krzywa jest monotoniczna, więc jednej wartości OCV odpowiada jeden SoC – ale relacja obowiązuje wyłącznie po relaksacji, nie pod obciążeniem

w równowadze; napięcie zaciskowe to to, co woltomierz pokazuje tu i teraz – najczęściej pod obciążeniem i przed relaksacją. Różnica między nimi jest dokładnie powodem, dla którego „pomiar SoC woltomierzem” zawodzi.

2. Dlaczego napięcie kłamie

Napięcie, które mierzymy na zaciskach pracującego ogniwa, to OCV pomniejszone (lub powiększone – zależnie od kierunku prądu) o szereg składników. W ujęciu modelu zastępczego zapisujemy to następująco:

$$v(t) = OCV(z(t)) - i(t)R_0 - \sum_j v_{C_j}(t)$$

Drugi składnik to omowy spadek napięcia na rezystancji wewnętrznej R_0 , proporcjonalny do prądu i pojawiający się natychmiast. Kolejne składniki v_{C_j} to napięcia na członach RC (ang. *resistor–capacitor*) modelujące nadpotencjały dyfuzyjne i polaryzację – narastają i zanikają z pewnymi stałymi czasowymi. Do tego dochodzi wpływ temperatury oraz histereza zależna od kierunku ostatniego prądu.

Najlepiej widać to na odpowiedzi ogniwa na impuls prądu (**rysunek 2**). W chwili załączenia obciążenia napięcie skokowo spada o człon omowy, następnie powoli osiada w miarę narastania nadpotencjałów. Po zdjęciu obciążenia skokowo wraca część omowa, a potem – przez minuty, a bywa, że godziny – napięcie relaksuje w stronę

prawdziwego OCV. Dopiero ta wartość końcowa nadaje się do odczytu SoC z krzywej z rysunku 1.

Drugim, często niedocenianym sprawcą jest histereza. To samo ogniwo przy tym samym SoC ma nieco wyższe OCV, jeśli ostatnio było ładowane, a niższe – jeśli rozładowywane. W ogniwach LFP zjawisko to jest szczególnie wyraźne i samowsobnie wprowadzić błąd odczytu napięciowego rzędu kilku–kilkunastu procent SoC.

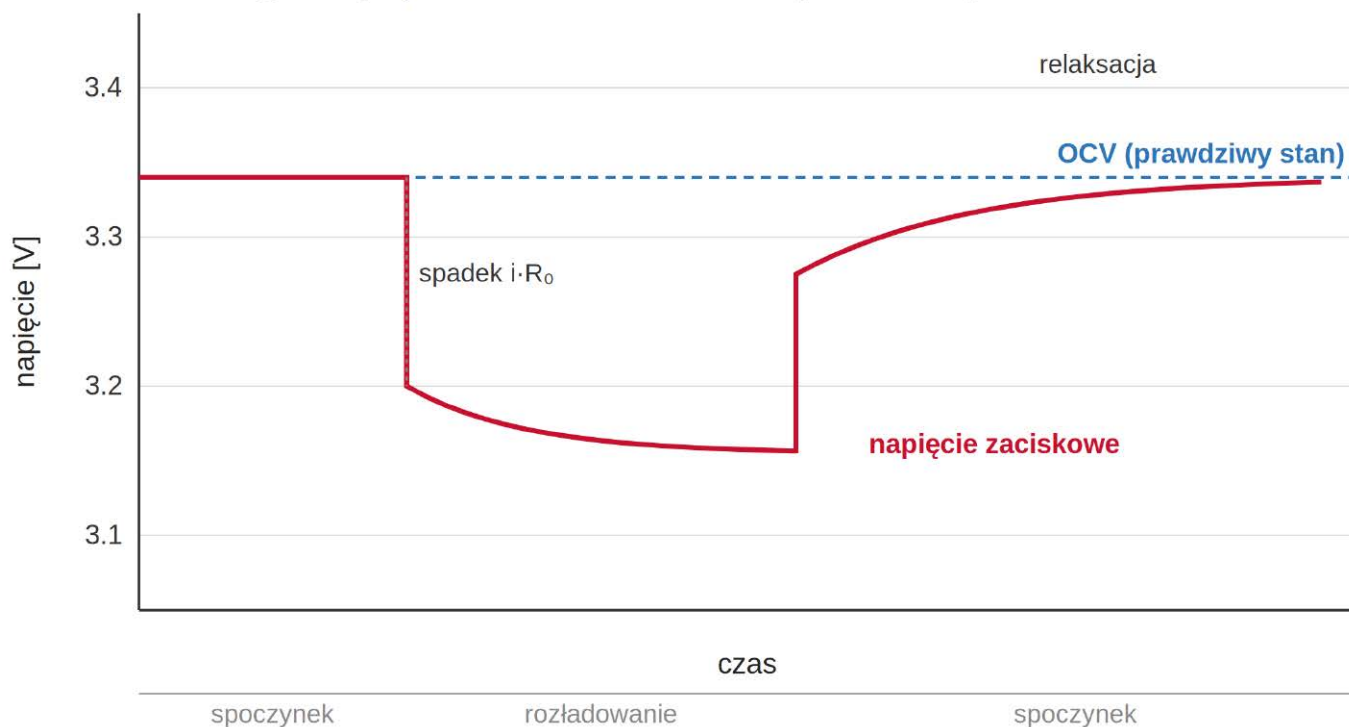
Praktyczny wniosek z tego punktu jest następujący: odczyt napięciowy ma sens dopiero po długim odpoczynku ogniwa, a w realnym urządzeniu – gdzie prąd płynie niemal stale – takiego luksusu zwykle nie ma. Woltomierz pokazuje stan elektryczny końcówek, a nie stan naładowania.

3. Przypadek LFP – płaskie plateau

Jeśli powyższe problemy dotyczą każdej chemii, to LFP doprowadza je do skrajności. Krzywa OCV–SoC ogniwa litowo-żelazowo-fosforanowego ma rozległe, niemal płaskie plateau (**rysunek 3**). W zakresie roboczym od ok. 20% do 80% SoC napięcie zmienia się o około 0,09 V – to mniej niż poziom szumów wielu typowych torów pomiarowych. Dla porównania chemia NMC (ang. *Nickel Manganese Cobalt*) ma w tym samym zakresie nachylenie wielokrotnie większe, dzięki czemu napięcie niesie informację o SoC [5, 6].

Konsekwencje są czysto praktyczne. Po pierwsze, na plateau odczyt OCV daje wysoce niepewny SoC – różnica

Dlaczego napięcie zaciskowe kłamie pod obciążeniem

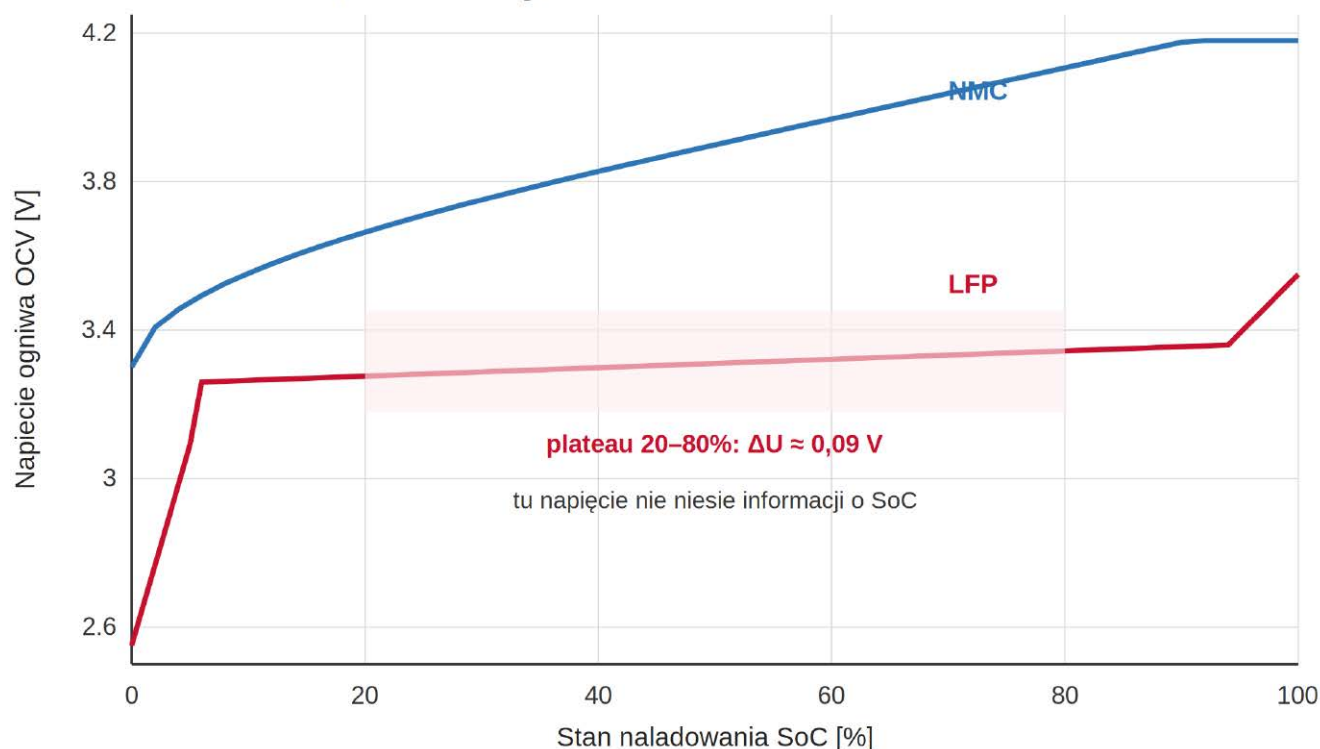


2. Odpowiedź napięcia zaciskowego na impuls prądu. Skok pionowy to omowy spadek na R_0 ; powolne osiadanie i powrotna relaksacja wynikają z nadpotencjałów. Prawdziwy SoC odpowiada dopiero zrelaksowanemu OCV (linia przerywana)

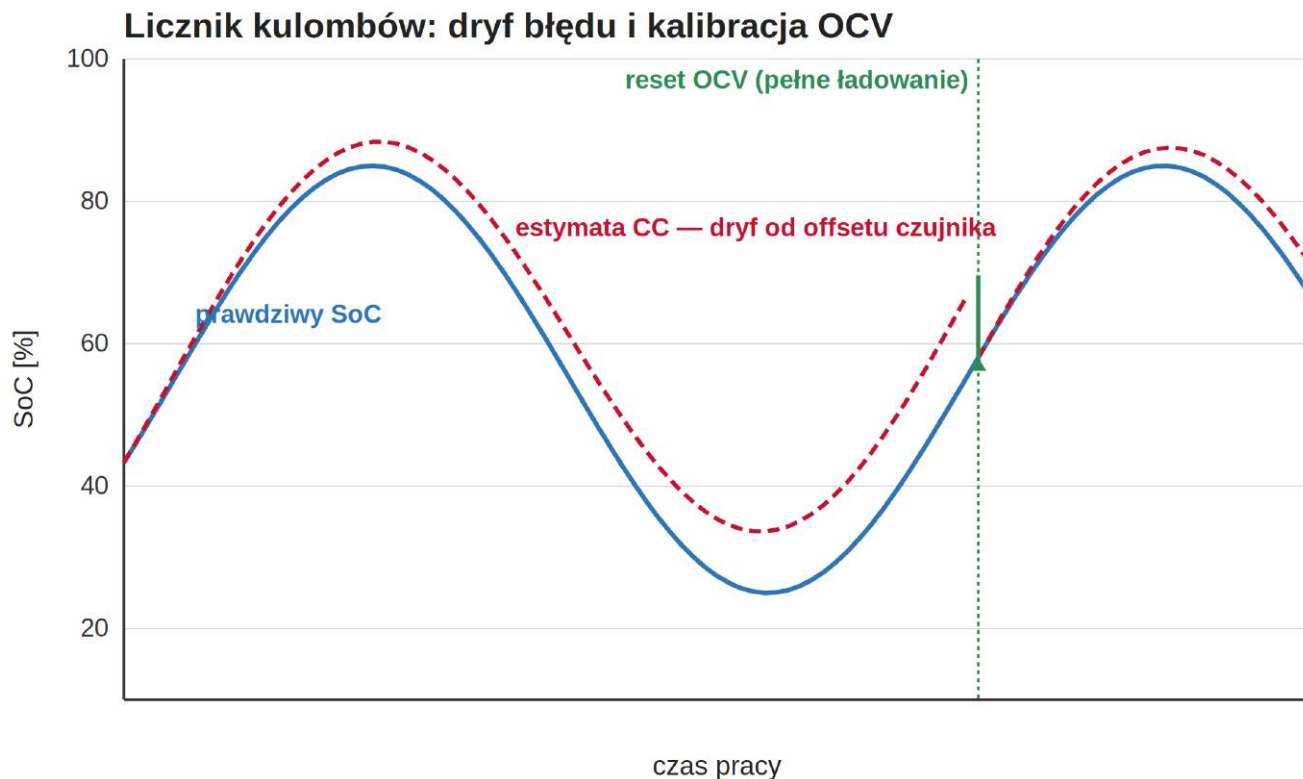
napięć między 30% a 70% SoC bywa mniejsza niż błąd pomiaru. Po drugie, estymator oparty wyłącznie na napięciu zachowuje się skokowo: trzyma niemal stałą wartość przez całe plateau, a potem – gdy ogniwo zejdzie z plateau przy

krańcach – „przeskakuje”, np. z 60% do 20% niemal natychmiast. W systemie objawia się to fałszywymi alarmami, przedwczesnymi odcięciami i myłymi prognozami stanu [6].

Dwie chemie, dwa światy: NMC vs LFP



3. Porównanie krzywych OCV-SoC: nachylona charakterystyka NMC i płaskie plateau LFP. W zacięzionym pasie 20...80% napięcie praktycznie nie zmienia się wraz ze stanem naładowania



4. Estymata licznika kulombów dryfuje od prawdziwego SoC wskutek offsetu czujnika prądu. Kalibracja OCV przy pełnym naładowaniu (reset) ściąga estymatę z powrotem do wartości prawdziwej; między resetami błąd narasta ponownie

Po trzecie, stąd bierze się popularne nieporozumienie: „13,2 V na baterii 12 V oznacza pełne naładowanie”. W rzeczywistości dla LFP zrelaksowane 13,2 V odpowiada zwykle 50...60% SoC – prawdziwie pełne ogniwo manifestuje się napięciem dopiero przy samej krawędzi plateau.

Stąd wniosek: dla LFP minimum akceptowalne to licznik kulombów z okresową kalibracją OCV w punktach skrajnych (pełne/puste), a docelowo – estymacja modelowa, która łączy prąd i napięcie. Tymi metodami zajmiemy się dalej (licznik kulombów i jego dryf oraz filtr Kalmana według Pletta).

4. Licznik kulombów i jego dryf

Skoro napięcie nie wystarcza, narzuca się pomysł odwrotny: zamiast zgadywać SoC z napięcia, liczyć bezpośrednio ładunek wpływający i wypływający z ogniwa. To metoda licznika kulombów (ang. *Coulomb counting*, CC) – w praktyce całkowanie prądu w czasie. W postaci dyskretnej, próbka po próbce, zapisujemy ją tak:

$$z_k = z_{k-1} - \frac{\eta_k \Delta t}{Q}$$

gdzie Δt to okres próbkowania. Zaleta jest zasadnicza: CC nie zależy od kształtu krzywej OCV, więc działa nawet na płaskim plateau LFP, które rozbroiło metodę napięciową z pkt. 3. W krótkim horyzoncie czasowym licznik kulombów potrafi być bardzo dokładny.

Problem jest jednak równie zasadniczy: licznik kulombów to estymator otwarty – bez sprzężenia zwrotnego,

które korygowałoby pomyłki. Każdy błąd raz popełniony zostaje w wyniku i się kumuluje. Źródła błędów są trzy.

- **Offset i dryf czujnika prądu.** Nawet niewielki, stały błąd zera b całkuje się w czasie do błędów SoC rosnącego liniowo:

$$\varepsilon_z(t) = \frac{1}{Q} \int_0^t b d\tau = \frac{bt}{Q}$$

Po godzinach pracy z offsetem rzędu miliamperów estymata potrafi rozjechać się o kilkanaście procent (rysunek 4).

- **Błąd warunku początkowego.** Całkowanie wymaga znajomości punktu startowego $z(t_0)$; jeśli nie znamy go dokładnie, błąd przenosi się na całą dalszą trajektorię.
- **Niepewność pojemności i sprawności.** Wartość Q nie jest stała – maleje wraz ze starzeniem (wrócimy do tego w pkt. 7), a sprawność kulombowa η oraz efekt Peukerta (zależność użytecznej pojemności od prądu) wprowadzają dodatkowe odchyłki.

Lekarstwem jest domknięcie pętli – okresowa kalibracja (reset) estymaty w punktach, gdzie napięcie znowu niesie informację. To krańce zakresu: blisko pełnego i pustego krzywa OCV–SoC jest stroma (por. rysunek 3), więc tam odczyt OCV daje wiarygodny SoC, którym można „przykleić” licznik z powrotem do prawdy (rysunek 4).

Mamy więc dwie metody o komplementarnych wadach: napięcie/OCV jest wiarygodne w równowadze i przy

Uwaga praktyczna – pułapka „top-up charging”

W typowym użytkowaniu LFP (magazyny energii, instalacje off-grid) bateria bywa trzymana stale w wąskim oknie, np. 30...80%, bez pełnego ładowania ani głębokiego rozładowania. Jest to korzystne dla trwałości ogni, ale pozbawia BMS punktów odniesienia: licznik kulombów nigdy nie trafia na stromy fragment krzywej, więc nie ma jak się zresetować – a dryf narasta bez korekty.

Praktyczna rada: zaplanować okresowe pełne ładowanie (rekalibrację), nawet jeśli aplikacja sama w sobie go nie wymaga.

krańcach, lecz bezradne na plateau i pod obciążeniem; licznik kulombów jest dobry w krótkim horyzoncie i na plateau, lecz dryfuje w długim. Naturalny wniosek inżynierski: połączyć je w jednym estymatorze, który samoczynnie waży zaufanie do każdego źródła. Tym estymatorem jest filtr Kalmana.

5. Estymacja modelowa – filtr Kalmana według Pletta

Przełom przyszedł w 2004 roku, gdy Gregory L. Plett w trzyczęściowej pracy w „Journal of Power Sources” pokazał, jak zastosować rozszerzony filtr Kalmana (ang. *Extended Kalman Filter*, EKF) do estymacji SoC pakietów litowo-jonowych [2]. Idea: nie wybierać między licznikiem kulombów a napięciem, lecz oprzeć estymację na modelu ogniwa i pozwolić filtrowi połączyć oba źródła informacji.

Formalnie filtr operuje na wektorze stanu, którego jedną ze składowych jest SoC. Model opisują dwa równania – dynamiki stanu i wyjścia:

$$x_k = f(x_{k-1}, u_{k-1}) + w_{k-1}$$

$$y_k = g(x_k, u_k) + v_k$$

Pierwsze równanie (predykcja) jest w istocie zegarem całkującym prąd: to licznik kulombów zaszyty w równaniu stanu. Drugie wiąże stan z mierzalnym

napięciem zaciskowym przez model zastępczy ECM (równanie z pkt. 2).

Filtr działa w pętli predykcja–korekcja (rysunek 5). W kroku predykcji model przewiduje nowy stan oraz napięcie wyjściowe. W kroku korekcji mierzone napięcie porównujemy z przewidzianym; różnicę – zwaną innowacją – mnożymy przez wzmocnienie Kalmana i tą poprawką korygujemy estymatę stanu:

$$\hat{x}_k^+ = \hat{x}_k^- + L_k \left(y_k - \hat{y}_k^- \right)$$

Wzmocnienie L_k nie jest stałe – filtr wylicza je z bieżących niepewności, decydując na bieżąco, na ile zaufać pomiarami, a na ile modelowi.

Tu tkwi elegancja rozwiązania. Na plateau LFP czułość napięcia na SoC jest znikoma, więc innowacja napięciowa niesie mało informacji – filtr automatycznie obniża jej wagę i bardziej ufa całkowaniu prądu. Przy krańcach, gdzie OCV jest strome, jest odwrotnie: napięcie mocno koryguje dryf licznika. To samoczynne ważenie rozwiązuje jednocześnie problem płaskiego plateau (pkt. 3) i dryfu licznika (pkt. 4) – bez ręcznych progów.

Dlaczego model zastępczy, a nie pełna elektrochemia? Jak zauważają Plett i Trimboli, ECM bardzo dobrze przewiduje napięcie ogniwa, więc oparta na nim estymacja SoC i energii jest dokładna i tania obliczeniowo [4]. Cena: model empiryczny nie odtwarza wewnętrznych zmiennych elektrochemicznych, więc stan zdrowia (ang. *State of Health*, SoH) i mechanizmy degradacji wymagają osobnego podejścia (do czego wrócimy w pkt. 7).

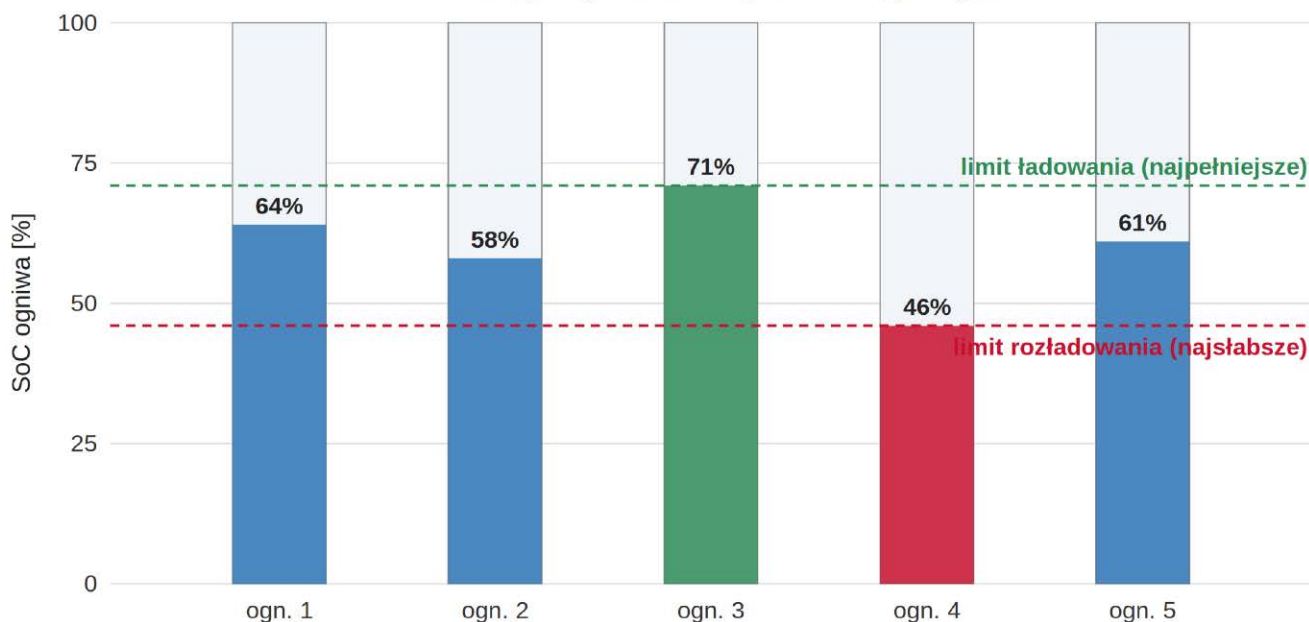
Filtr Kalmana nie jest jednak darmowy. Wymaga sparametryzowanego modelu: trzeba zidentyfikować rezystancje i stałe czasowe członów RC, krzywą OCV–SoC oraz ich zależność od temperatury i stopnia naładowania. Jakość estymaty jest dokładnie tak dobra, jak model. He, Xiong i Guo pokazali, że parametry modelu i SoC ogniwa LFP można estymować online, w trakcie pracy [7]; Baronti

Estymator modelowy SoC — rozszerzony filtr Kalmana



5. Estymator SoC oparty na rozszerzonym filtrze Kalmana. Krok predykcji (model ECM + całkowanie prądu) i krok korekcji (innowacja napięcia ważona wzmocnieniem Kalmana) działają w pętli ze sprzężeniem zwrotnym

Pakiet szeregowy: limit najsłabszego ogniwa



pojemność użyteczna pakietu = okno między liniami — mniejsze niż każdego ogniwa z osobna

6. W pakiecie szeregowym rozładowanie kończy się, gdy najsłabsze ogniwo osiągnie dolną granicę, a ładowanie – gdy najpełniejsze osiągnie górną. Pojemność użyteczna to zawężone okno między tymi liniami

i współpracownicy zademonstrowali podejście model-based łączące licznik ładunku z modelem napięcia dla ogniwa wysokoprądowych [8].

Wszystko, o czym dotąd mówiliśmy, dotyczyło pojedynczego ogniwa. Pakiet to osobna historia: ogniwa się rozjeżdżają, a użyteczna pojemność całości bywa zakładnikiem najsłabszej celi. Tym – oraz wpływem starzenia na samą definicję „100%” – zajmiemy się w pkt. 6 i 7.

6. Pakiet to nie ogniwo – rozjazd i balansowanie

Dotąd mówiliśmy o pojedynczym ogniwie. Realny akumulator to szereg (często z równoległymi gałęziami) wielu ogniw – a tu pojawia się nowy problem: ogniwa nie są identyczne i z czasem coraz bardziej się różnią. Davide Andrea, konstruktor BMS i autor monografii o systemach BMS dla dużych pakietów litowo-jonowych, formułuje regułę porządkującą projektowanie pakietów: użyteczna pojemność całości jest ograniczona przez najsłabsze ogniwo [9].

Mechanizm jest prosty (rysunek 6). Ogniwa połączone szeregowo przewodzą ten sam prąd, ale mają nieco różne SoC i pojemności. Przy rozładowaniu jako pierwsze osiągnie dolną granicę ogniwo najsłabsze – i w tym momencie BMS musi przerwać, choć pozostałe mają jeszcze zapas. Przy ładowaniu pierwsze nasyci się ogniwo najpełniejsze – i znowu trzeba przerwać. Okno między tymi progami, czyli faktyczna pojemność użyteczna pakietu, jest węższe niż pojemność któregośkolwiek ogniwa z osobna.

Skąd biorą się różnice? Z rozrzutu produkcyjnego pojemności i rezystancji, z nierównomierniej

temperatury w pakiecie (ogniwa w środku grzeją się bardziej), z różnic w prądzie samorozładowania oraz z nierównego starzenia. Te drobne odchyłki kumulują się cykl po cyklu – pakiet „rozjeżdża się”.

Lekarstwem jest balansowanie ogniw (rysunek 7). W wariantcie pasywnym nadmiar energii z najpełniejszych ogniw jest rozpraszany w rezystorach upustowych, aż wszystkie zrównają się z najsłabszym – proste i tanie, lecz energia idzie w ciepło, a tempo jest ograniczone. W wariantcie aktywnym ładunek jest przenoszony z ogniw pełniejszych do słabszych przez przetwornice lub kondensatory – wydajniej i szybciej, kosztem złożoności i ceny [12].

Wniosek dla estymacji: SoC pakietu nie jest średnią SoC ogniw. Dobrze zaprojektowany BMS śledzi stan każdej celi (lub grupy równoległej) z osobna i raportuje SoC całości względem ograniczeń najsłabszego i najpełniejszego ogniwa. Estymator z pkt. 5 działa wtedy nie na jednym, lecz na wielu ogniwach, a balansowanie utrzymuje okno użyteczne możliwie szerokie.

7. SoC dryfuje także przez starzenie

Jest jeszcze jeden, podstępny powód, dla którego „stan naładowania” wymyka się prostym definicjom: punkt odniesienia sam się zmienia. SoC liczymy względem pojemności pełnego ogniwa – ale ta pojemność maleje przez całe życie akumulatora. To, co dziś jest „100%”, za dwa lata będzie odpowiadało mniejszej liczbie amperogodzin.

Ten ubytek opisuje stan zdrowia (SoH), zdefiniowany jako stosunek bieżącej pojemności do znamionowej:

Balansowanie ogniw: pasywne i aktywne



7. Dwa podejścia do balansowania. Pasywne rozprasza nadmiar energii jako ciepło; aktywne przenosi ładunek między ogniwami z niewielkimi stratami

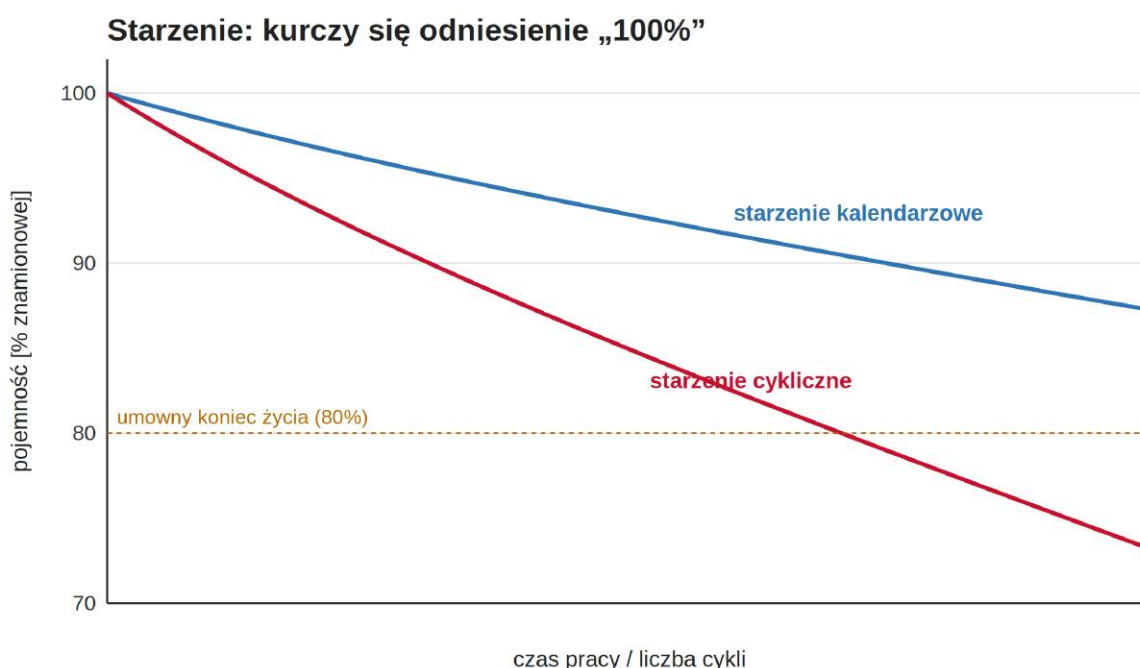
$$SoH = \frac{Q_{aktualne}}{Q_{znamionowe}} \cdot 100\%$$

SoC i SoH są sprzężone: estymator liczący SoC względem nieaktualnej, zawyżonej pojemności będzie systematycznie się mylił. Dlatego dojrzałe BMS estymują obie wielkości i okresowo aktualizują Q.

Degradacja ma dwa składniki, które zespół Andreasa Jossena (Politechnika Monachijska) rozdzielił w szeroko cytowanych badaniach [10]. Starzenie kalendarzowe zachodzi z upływem czasu, niezależnie od użycia – zależy głównie od temperatury i poziomu naładowania, w jakim ogniwo jest przechowywane. Starzenie cykliczne zależy od przepuszczonego ładunku: liczby i głębokości cykli (rysunek 8). W realnym urządzeniu oba składniki się nakładają.

Czynniki przyspieszające degradację są dobrze rozpoznane: wysoka temperatura, przechowywanie przy wysokim SoC, głębokie cykle i wysokie prądy ładowania. Spotnitz i Franklin pokazali, jak wysokie prądy i temperatury pogłębiają degradację, a w skrajnych warunkach prowadzą do niestabilności [11]. Stąd praktyczne zalecenia: unikać długiego trzymania ogniw przy 100% i w wysokiej temperaturze, a szybkie ładowanie traktować jako kompromis między wygodą a żywotnością.

Dla projektanta oznacza to, że estymator SoC nie może być „ustawiony i zapomniany”. Pojemność, rezystancja i kształt krzywej OCV–SoC dryfują z wiekiem; model trzeba albo adaptować online (jak w pracach He i Xionga, [7]), albo okresowo rekalkulować.



8. Starzenie kalendarzowe i cykliczne obniżają pojemność ogniwa. Ponieważ SoC odnosimy do pełnej pojemności, malejące „100%” oznacza, że ten sam odczyt SoC odpowiada z czasem coraz mniejszej energii

8. W pracowni – jak zrobić to dobrze

Zbierzmy wnioski w konkretne wskazówki warsztatowe.

- **Wyznacz własną krzywą OCV–SoC ogniwa.** Nie ufaj tabelom z internetu – rozrzut między partiami bywa istotny. Rozładuj i naładuj ogniwo małym prądem ($C/20$ – $C/30$) albo metodą impulsową z długimi przerwami na relaksację, mierząc OCV po ustabilizowaniu; stabelaryzuj osobno dla ładowania i rozładowania (histereza).
- **Dobierz tor pomiaru prądu pod licznik kulombów.** Bocznik lub przetwornik o małym offsecie i niskim dryfie temperaturowym, z przetwornikiem ADC o adekwatnej rozdzielczości. Pamiętaj: zabójczy jest offset, nie szum – szum uśrednia się w całość, offset się kumuluje (pkt. 4).
- **Zaplanuj resety OCV.** Zdefiniuj warunki kalibracji estymaty: po dłuższym spoczynku przy znanej temperaturze, najlepiej blisko krańców zakresu, gdzie OCV jest strome. Dla LFP wymuś okresowe pełne ładowanie (pkt. 4 – pułapka top-up).
- **Nie diagnozuj LFP samym napięciem.** Na plateau różnica napięć jest poniżej szumu – alarmy i prognozy

Co zapamiętać

- Napięcie to nie SoC; SoC nie jest mierzalny, lecz estymowalny.
- OCV mapuje na SoC tylko w równowadze; pod obciążeniem napięcie kłamie.
- LFP: płaskie plateau czyni metodę napięciową bezużyteczną w środku zakresu.
- Licznik kulombów jest dokładny krótkoterminowo, lecz dryfuje – wymaga resetów OCV.
- Filtr Kalmana (EKF) łączy oba źródła i samoczynnie waży zaufanie do każdego.
- Pakiet ogranicza najniższe ogniwo; trzeba balansować i estymować per celi.
- Pojemność maleje z wiekiem – odniesienie „100%” się kurczy (SoH).

Checklista projektanta BMS

- ✓ Zmierzona własna krzywa OCV–SoC (ładowanie/rozładowanie, zależność od temperatury).
- ✓ Tor prądowy o niskim offsecie i dryfie; znana rozdzielczość ADC.
- ✓ Zdefiniowane warunki resetu OCV (spoczynek, krańce zakresu).
- ✓ Estymator modelowy (EKF) zamiast progów napięciowych – zwłaszcza dla LFP.
- ✓ Estymacja per ogniwo + balansowanie (pasywne lub aktywne).
- ✓ Bieżąca aktualizacja pojemności Q i SoH.
- ✓ Strategia okresowej pełnej kalibracji (przeciw pułapce top-up).

opieraj na liczniku kulombów skorygowanym estymatorem, nie na progach napięciowych (pkt. 3).

- **Estymuj per ogniwo, raportuj per pakiet.** Śledź najniższą i najpełniejszą celę; SoC pakietu odnoś do ich ograniczeń; utrzymuj balansowanie (pkt. 6).
- **Aktualizuj pojemność.** Przy każdej pełnej kalibracji odśwież oszacowanie Q i SoH; estymator liczący SoC względem starej pojemności dryfuje (pkt. 7).

WM, JS

Ekspertiści przywołani w artykule

Gregory L. Plett – profesor inżynierii elektrycznej i komputerowej na University of Colorado Colorado Springs (UCCS). Autor dwutomowej monografii „Battery Management Systems” (Artech House, 2015) i pionier estymacji SoC opartej na rozszerzonym filtrze Kalmana (seria prac z 2004 r.); współtwórca ogólnodostępnych kursów o algorytmach BMS.

Davide Andrea – inżynier i przedsiębiorca, twórca firmy Elithion (systemy BMS), autor monografii „Battery Management Systems for Large Lithium-Ion Battery Packs” (Artech House, 2010) – jednego z najczęściej przywoływanych praktycznych źródeł o projektowaniu pakietów i balansowaniu.

Andreas Jossen – profesor Politechniki Monachijskiej (TUM), kierujący Katedrą Elektrycznych Technologii Magazynowania Energii; jego zespół opublikował szeroko cytowane prace rozdziałające starzenie kalendarzowe i cykliczne ogniwo litowo-jonowych.

Bibliografia

- [1] G. L. Plett, Battery Management Systems, Vol. I: Battery Modeling (ISBN 978-1-63081-023-8) oraz Vol. II: Equivalent-Circuit Methods (ISBN 978-1-63081-027-6), Artech House, Norwood (MA), USA 2015. <http://mocha-java.uccs.edu/BMS1/>
- [2] G. L. Plett, „Extended Kalman filtering for battery management systems of LiPB-based HEV battery packs”, Journal of Power Sources, vol. 134, no. 2, 2004: cz. 1, s. 252–261 (doi:10.1016/j.jpowsour.2004.02.031); cz. 2, s. 262–276 (doi:10.1016/j.jpowsour.2004.02.032); cz. 3, s. 277–292 (doi:10.1016/j.jpowsour.2004.02.033). <http://mocha-java.uccs.edu/dossier/RESEARCH/2004jps1-.pdf>
- [3] C. H. Weng, J. Sun, H. Peng, „A unified open-circuit-voltage model of lithium-ion batteries for state-of-charge estimation and state-of-health monitoring”, Journal of Power Sources, vol. 258, 2014, s. 228–237 (doi:10.1016/j.jpowsour.2014.02.026). https://huei.engin.umich.edu/wp-content/uploads/sites/186/2015/02/2014_OCV_Model_JPS.pdf
- [4] „An Interview with Gregory L. Plett and M. Scott Trimboli”, Artech House Insider, 25 stycznia 2024. <https://blog.artechhouse.com/2024/01/25/an-interview-with-gregory-l-plett-and-m-scott-trimboli/>
- [5] PowerTech Systems, „Lithium-Ion State of Charge (SoC) measurement”, Tech Corner [dostęp: 2026]. <https://www.powertechsystems.eu/tech-corner/lithium-ion-state-of-charge-soc-measurement/>
- [6] InSum Energy, „SoC estimation methods: Coulomb Counting vs Kalman Filtering for LiFePO4 batteries” [dostęp: 2026]. <https://www.insumenergy.com/soc-estimation-methods-coulomb-counting-vs-kalman-filtering-for-lifepo4-batteries/>
- [7] H. He, R. Xiong, H. Guo, „Online estimation of model parameters and state-of-charge of LiFePO4 batteries in electric vehicles”, Applied Energy, vol. 89, no. 1, 2012, s. 413–420.
- [8] A. Baronti, G. Fantechi, R. Roncella, R. Saletti, „State-of-charge estimation for high-power Li-ion cells using a model-based approach”, IEEE Transactions on Instrumentation and Measurement, vol. 60, no. 9, 2011, s. 3275–3285.
- [9] D. Andrea, Battery Management Systems for Large Lithium-Ion Battery Packs, Artech House, Norwood (MA), USA 2010 (ISBN 978-1-60807-104-3). <https://ieeexplore.ieee.org/document/9100544/>
- [10] E. Schuster, P. Keil, S. F. Schuster, A. Jossen, „Calendar and cycling ageing of lithium-ion batteries”, Journal of Power Sources, vol. 279, 2015, s. 752–762.
- [11] R. Spotnitz, J. Franklin, „Abuse behavior of high-power, lithium-ion cells”, Journal of Power Sources, vol. 113, no. 1, 2003, s. 81–100.
- [12] C. R. S. Santos, A. M. A. C. Rocha, „Method and algorithm for efficient cell balancing in the lithium-ion battery”, CEUR Workshop Proceedings, vol. 3842, 2024, paper 16. <https://ceur-ws.org/Vol-3842/paper16.pdf>

θ_{JA} z noty katalogowej to fikcja

W notcie katalogowej każdego elementu mocy jest ta jedna, wygodna liczba – opór cieplny: $\theta_{JA} = XY^{\circ}\text{C/W}$. Kusi, żeby wstawić ją do wzoru $T_J = T_A + P \cdot \theta_{JA}$, odczytać temperaturę złącza i mieć spokój. Kłopot w tym, że ta liczba opisuje znormalizowaną płytkę pomiarową w nieruchomym powietrzu – a nie Twoją płytkę i Twoje otoczenie. Co więcej, sami producenci to zastrzegają: nota ROHM wprost pisze, że θ_{JA} służy wyłącznie do porównywania obudów w znormalizowanym środowisku i nie nadaje się do przewidywania temperatury w konkretnej aplikacji. Pokażemy, skąd bierze się ta liczba, dlaczego myli się o dziesiątki stopni i czym ją zastąpić, opierając się na notach ROHM, onsemi i Texas Instruments (TI) oraz na klasycznych pracach o termice obudów.

Skąd bierze się jedna liczba

Zacznijmy od faktu, który w codziennej pracy łatwo przeoczyć: **θ_{JA} nie jest liczone – jest mierzone**. Producent odczytuje je z eksperymentu prowadzonego według norm JEDEC (Joint Electron Device Engineering Council). Zgodnie z definicją przywołaną w notcie aplikacyjnej ROHM [7], θ_{JA} to opór cieplny od struktury czynnej półprzewodnika do otoczenia o konwekcji naturalnej (nieruchome powietrze); mierzy się je na płytce zgodnej z JESD51-3, -5 i -7, w środowisku wg JESD51-2A, a temperaturę złącza wyznacza elektrycznie – z napięcia przewodzenia złącza p-n w MOSFET (tranzystor polowy z izolowaną

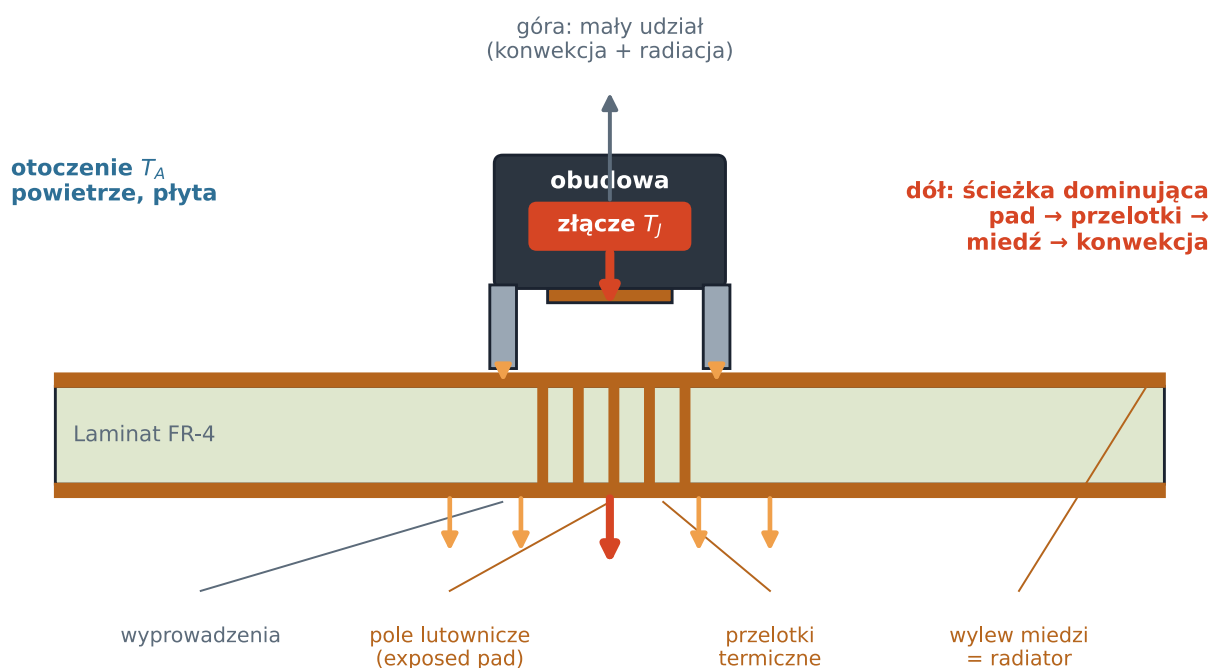
bramką), użytego jako czujnik, skalibrowany współczynnikiem K [1] [7]. Definicja jest wtedy trywialna:

$$\theta_{JA} = \frac{T_J - T_A}{P}$$

Standardowa płytkę ma wymiary 76,2 × 114,3 mm i grubość 1,6 mm; całość zamyka się w szczelnej komorze z materiału o małej przewodności cieplnej, by odciąć ruch powietrza [2] [7]. I tu pada zdanie, które warto powiesić nad biurkiem – nota ROHM stwierdza wprost, że celem pomiaru θ_{JA} jest jedynie porównywanie obudów w znormalizowanym środowisku, a nie przewidywanie ich zachowania w konkretnej aplikacji [7]. **Zmień płytkę, ilość miedzi**

Rzeczywiste ścieżki odprowadzania ciepła z obudowy SMD

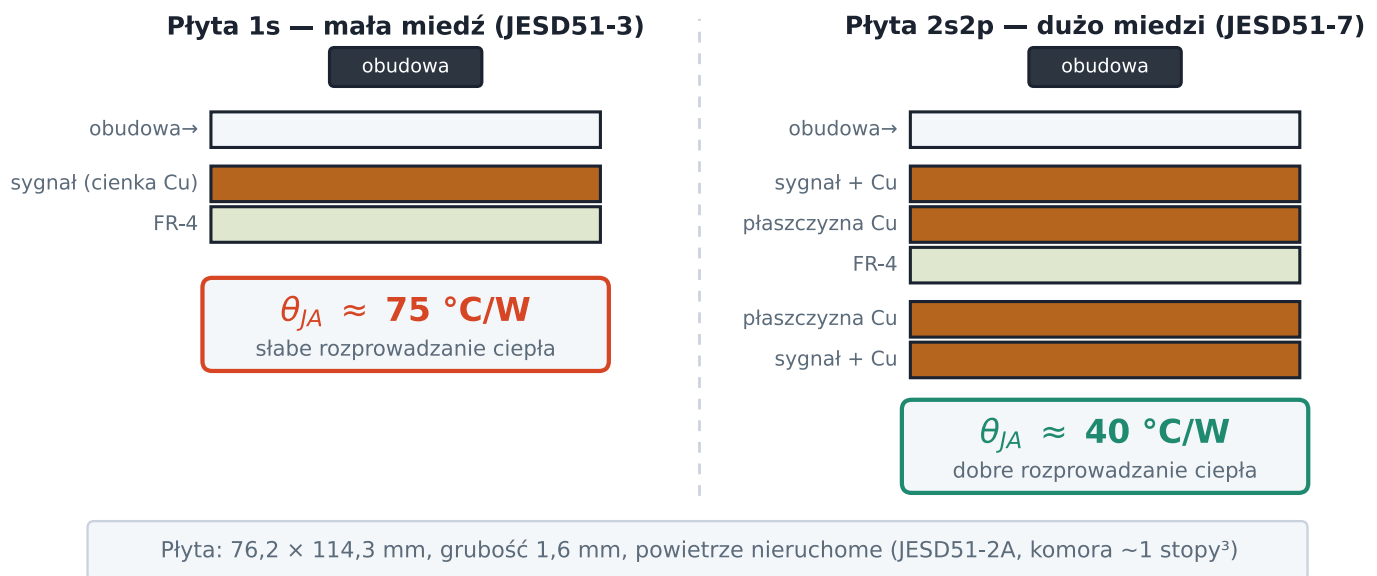
θ_{JA} nie jest cechą samej obudowy – obejmuje płytkę i otoczenie



1. W obudowie SMD większość ciepła odpływa w dół, do płytki (potwierdza to nota ROHM). To płytkę jest radiatorem – dlatego θ_{JA} zależy od niej, a nie tylko od obudowy

Ta sama obudowa, dwie znormalizowane płyty JEDEC

Wartość θ_{JA} zależy od płyty, na której ją zmierzono



2. Ta sama obudowa, dwie płyty JEDEC. Płyta 1s i 2s2p dają istotnie różne θ_{JA} – najlepszy dowód, że parametr opisuje układ, nie element

albo przepływ powietrza – a θ_{JA} się zmieni. Dlatego w realnym układzie ta liczba jest fikcją.

Dlaczego to nie jest właściwość obudowy

Największe nieporozumienie polega na traktowaniu θ_{JA} jak parametru obudowy – tak jak traktujemy rezystancję rezystora. Tymczasem **θ_{JA} to parametr całego systemu**: obudowa plus płytka plus otoczenie. Klasyk zagadnienia, Clemens Lasance z Philips Research, w pracy o operacyjnym oporze cieplnym obudów [5] pokazał to już ćwierć wieku temu: θ_{JA} zależy od warunków brzegowych i nie jest wielkością wewnętrzną elementu. To on rozpropagował ideę modeli termicznych niezależnych od warunków brzegowych (BCI, ang. *boundary-condition independent*) – rozwiniętą w europejskich projektach DELPHI i PROFIT – czyli bezpośrednią odpowiedź na ułomność jednej liczby. Współczesne noty mówią to samo: onsemi w karcie sterownika FAN3121 opatruje wszystkie parametry cieplne przypisem, że są estymatami z symulacji i zależą od aplikacji [9].

Żeby to zobaczyć, prześledźmy drogę ciepła (**rysunek 1**). Złącze oddaje ciepło trzema drogami: w górę przez wierzch obudowy, na boki przez wyprowadzenia i w dół przez pole lutownicze (*exposed pad*), przelotki i wylew miedzi. W obudowach SMD (ang. *Surface-Mount Device* – montaż powierzchniowy) – QFN (ang. *Quad Flat No-leads*), DPAK (obudowa mocy TO-252), SOIC (ang. *Small Outline Integrated Circuit*) z padem – **droga dominująca prowadzi w dół, do płytki**. Nota ROHM ujmuje to jednym zdaniem: element SMD oddaje większość ciepła do PCB (ang. *Printed*

Circuit Board – płytka drukowana) [7]. To warstwa miedzi, a nie obudowa, jest rzeczywistym radiatorem.

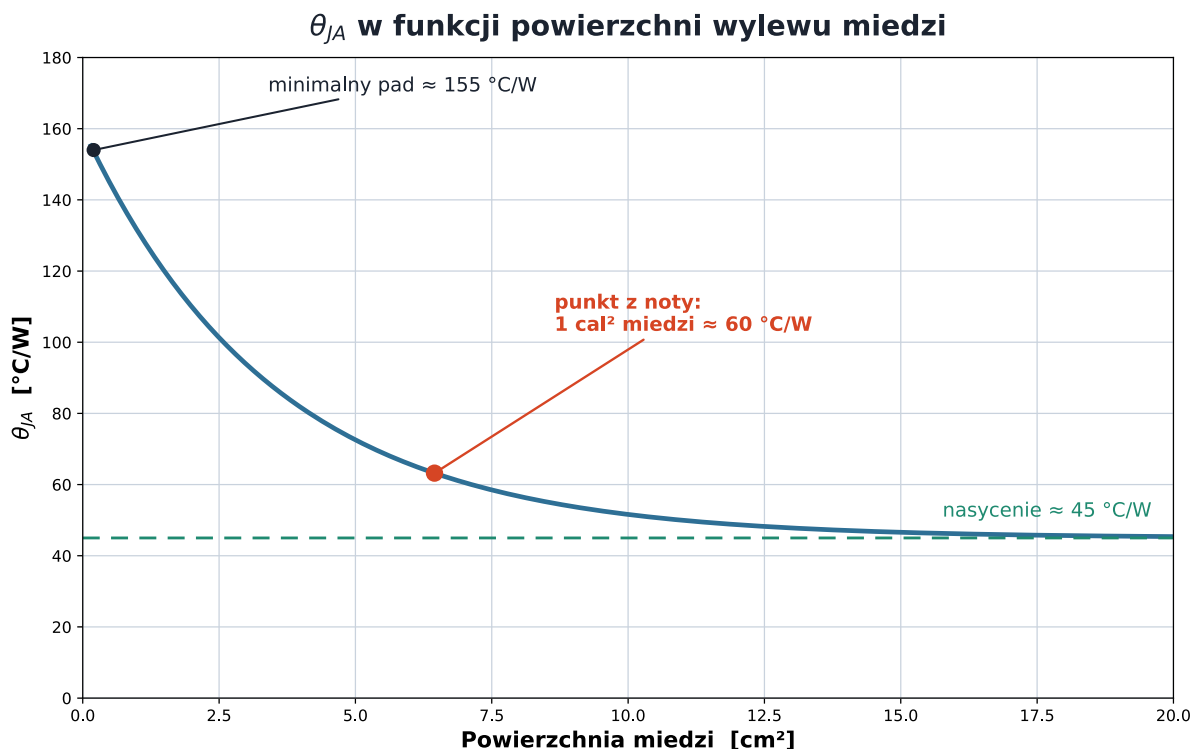
Wniosek jest oczywisty: mierząc θ_{JA} na **płytkce JEDEC**, mierzymy w dużej mierze płytkę JEDEC, a nie sam element. Parametrem obudowy jest dopiero θ_{JC} (złącze–obudowa), o którym za chwilę.

Dwie płyty, dwie prawdy – i dwie obudowy

Najprostszy dowód, że θ_{JA} nie należy do obudowy, to zestawienie dwóch znormalizowanych płyt (**rysunek 2**). Ta sama obudowa na płycie 1s (mało miedzi) i na płycie 2s2p (dużo miedzi) potrafi wykazać θ_{JA} różniące się niemal dwukrotnie. Efekt widać nawet wewnątrz jednej noty: w karcie FAN3121 ta sama struktura w obudowie SOIC-8 ma $\theta_{JA} \approx 87 \text{ °C/W}$, a w wersji SOIC-8 z padem (EP) $\approx 75 \text{ °C/W}$ – różnicę robi samo pole lutownicze [9].

onsemi nie zostawia tu złudzeń: w przypisie do θ_{JA} zapisuje, że wartość **zależy od projektu PCB, chłodzenia i przepływu powietrza**, a podana liczba dotyczy konwekcji naturalnej bez radiatora, na płytce 2S2P wg JESD51-2/-5/-7 [9]. Bywa i tak, że producent zmienia bazę odniesienia – w dokumentacji tranzystora mocy GaN (azotek galu) LMG3410 Texas Instruments zrewidowało θ_{JA} do warunków standardowej płytki JEDEC [13], co samo w sobie dowodzi, że liczba jest przypisana do założonej płytki. A ponieważ jeden parametr nie wystarcza, noty podają dziś zestaw – OmniVision w karcie WS4666D umieszcza obok siebie $R\theta_{JA}$ oraz $R\theta_{JC}(\text{top})$ i $R\theta_{JC}(\text{bot})$ [14].

θ_{JA} bez przypisu o płytce i przepływie nie znaczy nic. To pierwszy nawyk, który trzeba wyrobić: zanim użyjesz



3. θ_{JA} w funkcji powierzchni warstwy miedzi – kształt zgodny z krzywą ROHM dla HTSOP-J8. Popularny punkt „1 cal²” leży w stromej części: mała zmiana powierzchni miedzi to duża zmiana θ_{JA}

θ_{JA} , znajdź przypis. Jeśli go nie ma, potraktuj liczbę jako orientacyjną i nic więcej.

Płytką jest radiator

Skoro ciepło ucieka głównie w miedź, to **powierzchnia warstwy miedzi jest projektową dźwignią**, którą masz w ręku. Zależność θ_{JA} od pola miedzi ma charakterystyczny kształt nasycenia (rysunek 3): przy minimalnym padzie θ_{JA} jest bardzo wysokie, gwałtownie spada wraz z dodawaniem miedzi, po czym wypłaszcza się. Dokładnie taką krzywą pokazuje nota ROHM dla obudowy HTSOP-J8 na czterowarstwowej płycie JEDEC – z ważną uwagą: **wartość z noty leży na prawym, nasyconym końcu krzywej**, a różnicę na wykresie robi wyłącznie zmiana pola miedzi [7].

Nota ROHM wylicza przy tym, że poza samym polem miedzi na θ_{JA} wpływają grubość płytki, liczba warstw, grubość folii miedzianej, rozmieszczenie przelotek termicznych i pojemność komory [7] – czyli komplet parametrów, których w θ_{JA} „z nagłówka” nie ma. Praktyczne dźwignie to zatem: powierzchnia i grubość warstwy miedzi, przelotki łączące pad z płaszczyzną na spodzie oraz liczba warstw. Przy projektowaniu samych ścieżek i pól miedzi pomocne są wytyczne IPC-2152 (norma organizacji IPC – Association Connecting Electronics Industries) (i starsze IPC-2221) [15]. Płytką dwuwarstwową z minimalnym padem to termiczny dramat; czterowarstwową z płaszczyzną i gęstym polem przelotek – zupełnie inny świat.

Powietrze też się liczy

Drugim czynnikiem, którego nagłówkowe θ_{JA} nie uwzględnia, jest ruch powietrza. Wartość z noty mierzona jest w powietrzu nieruchomym – to najgorszy przypadek konwekcji. Skromny nawiew wyraźnie obniża θ_{JA} (rysunek 4); przy 2 m/s potrafi ono zmaleć niemal o połowę.

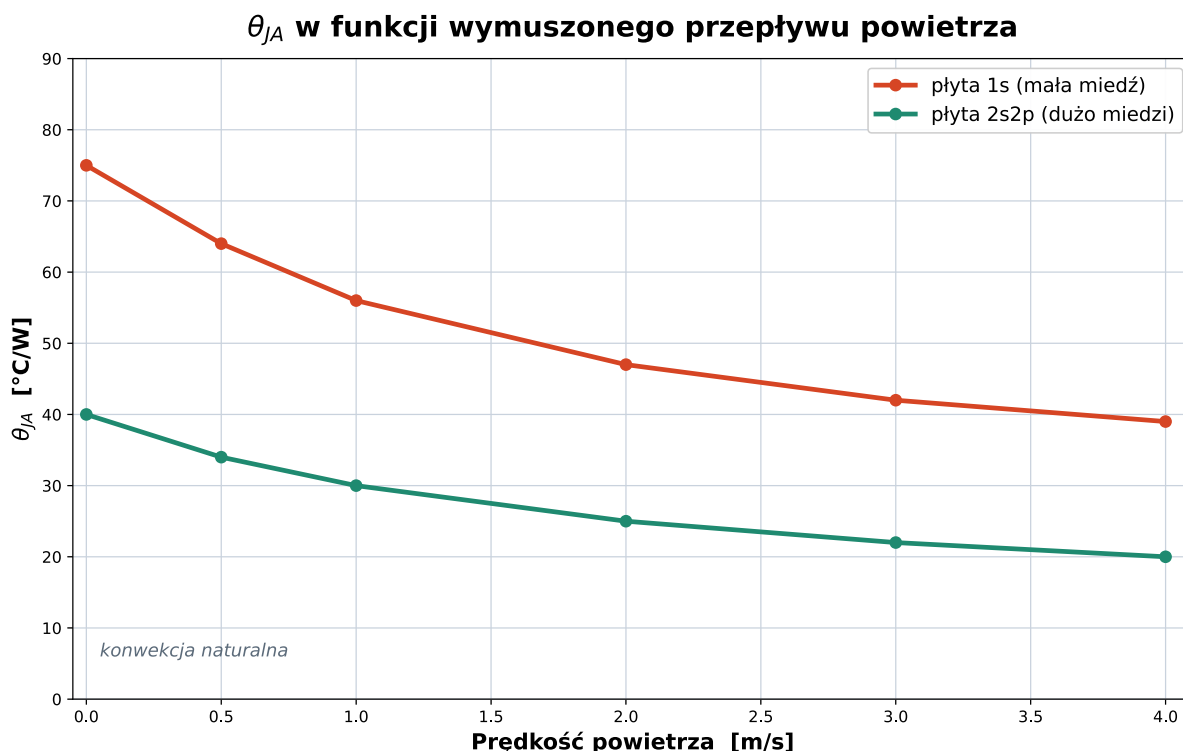
Dochodzą dwa efekty łatwe do przeoczenia. Po pierwsze **orientacja płytki**: przy konwekcji naturalnej płytka pionowa chłodzi się lepiej niż pozioma, bo wspomaga kominowy ruch powietrza. Po drugie **radiacja**: jej udział rośnie z czwartą potęgą temperatury bezwzględnej i zależy od emisyjności – przy dużych przyrostach temperatury bywa zauważalny, ale matowa obudowa promieniuje inaczej niż błyszcząca. Żaden z tych efektów nie jest zaszyty w jednej liczbie θ_{JA} .

Kiedy θ_{JA} jest jednak przydatne

Nie chodzi o to, że θ_{JA} jest bezwartościowe – chodzi o to, do czego wolno go używać. Nota ROHM podaje dwa dozwolone zastosowania [7]. Pierwsze to **ranking obudów**: niższe θ_{JA} (przy identycznych założeniach dotyczących płytki) oznacza lepsze odprowadzanie ciepła. Drugie to **względne oszacowanie**, o ile stopni różni się przyrost temperatury dwóch elementów:

$$\Delta T_J = (\theta_{JA2} - \theta_{JA1}) \cdot P$$

Do tego samego wniosku dochodzi Darwin Edwards z Texas Instruments – współautor norm płytek JEDEC i kanonicznej noty o metrykach cieplnych. Ujmuje on rzecz



4. θ_{JA} maleje z prędkością wymuszonego przepływu. Wartość z noty (przepływ zerowy) to przypadek najgorszy – realny nawiew działa na Twoją korzyść, ale nota o tym nie mówi

bez ogródek: θ_{JA} jest zarazem najczęściej raportowanym i najczęściej nadużywanym parametrem cieplnym, bo próbuje się nim szacować temperaturę złącza w układzie, do czego nie jest przeznaczony [11].

Aczgorobićnie wolno? ROHM nazywa to wprost nadużyciem θ_{JA} (misuse): wstawiania nagłówkowego θ_{JA} do wzoru z wiarą, że da on prawdziwą temperaturę złącza w Twoim urządzeniu [7]:

$$T_J = T_A + \theta_{JA} \cdot P$$

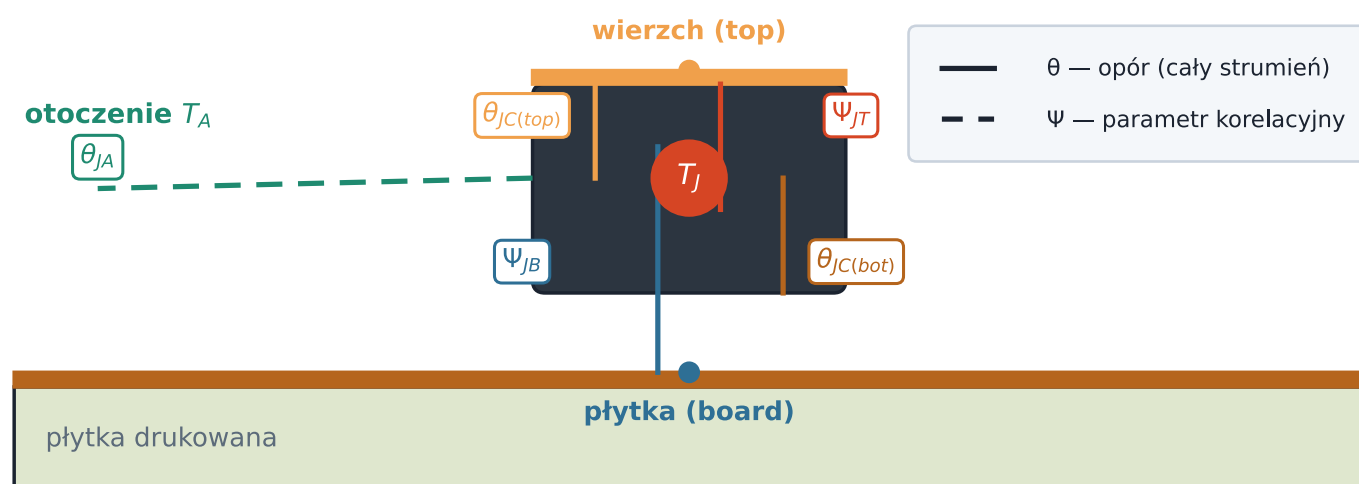
Ten wzór jest poprawny wyłącznie w warunkach, w których θ_{JA} zmierzono – na tamtej płytce, w tamtej komorze, w tamtym powietrzu. W realnym układzie rozdział ciepła jest inny, więc i θ_{JA} jest inne.

Lepsze narzędzia: θ_{JC}, ψ_{JT}, ψ_{JB}

Na szczęście noty podają dziś zestaw parametrów, które – użyte poprawnie – pozwalają wyznaczyć T_J rzetelnie. Ich punkty odniesienia pokazuje rysunek 5, a zestawienie – tabela 1.

Mapa parametrów cieplnych obudowy

Cztery punkty odniesienia: złącze, wierzch, płytka, otoczenie



5. Mapa parametrów cieplnych obudowy. θ (linie ciągłe) to prawdziwe opory mierzone przy niemal całym strumieniu ciepła jedną drogą; ψ (linie przerywane) to parametry korelacyjne

Tabela 1. Parametry cieplne obudowy – co znaczą i do czego służą

Parametr	Co opisuje	Mierzony przy	Do czego służy
θ_{JA}	złącze → otoczenie (obudowa + płytka + powietrze)	na płycie JEDEC, powietrze stoi	ranking obudów; NIE do przewidywania T_J
$\theta_{JC(top)}$	złącze → wierzch obudowy	zimna płyta na wierzchu, ~1D	projekt z radiatorem od góry
$\theta_{JC(bot)}$	złącze → pole lutownicze (spód)	zimna płyta pod padem, ~1D	projekt z radiatorem/płaszczyną od spodu
θ_{JL}	złącze → wyprowadzenia (i pad)	spód wyprowadzeń lutowanych do PCB	modele kompaktowe, symulacje
Ψ_{JT}	korelacja T_J z temperaturą wierzchu	na płycie JEDEC (rozdziel mocy)	wyznaczenie T_J z pomiaru wierzchu w realnym układzie
Ψ_{JB}	korelacja T_J z temperaturą płytki	na płycie JEDEC (rozdziel mocy)	wyznaczenie T_J z pomiaru płytki przy obudowie

θ_{JC} (złącze–obudowa) to prawdziwy opór cieplny i prawdziwy parametr obudowy. Mierzy się go, przykładając zimną płytę do obudowy tak, by wymusić przepływ niemal wyłącznie tą jedną drogą – stąd θ_{JC} jest niezależne od otoczenia. Dla obudów z padem rozróżnia się $\theta_{JC(top)}$ i $\theta_{JC(bot)}$; onsemi zaznacza, że $\theta_{JC(top)}$ zakłada utrzymanie wierzchu w stałej temperaturze przez radiator od góry [9]. Użyteczny jest ten od strony, którą realnie chłodzi.

Ψ_{JT} (złącze–wierzch) i Ψ_{JB} (złącze–płytki) to nie opory, lecz parametry korelacyjne (ang. *thermal characterization parameters*). Idea jest prosta: zmierz w swoim układzie realną temperaturę – wierzchu obudowy albo płytki tuż przy niej – i dodaj małą poprawkę [7] [9]:

$$T_J = T_{top} + \Psi_{JT} \cdot P$$

$$T_J = T_{board} + \Psi_{JB} \cdot P$$

Nota ROHM pokazuje, że przy zmianie pola miedzi Ψ_{JT} niemal się nie zmienia – w przeciwieństwie do θ_{JA} na tym samym wykresie [7]. To właśnie czyni je użytecznym w projektowaniu, o czym za chwilę.

θ kontra Ψ – różnica, którą trzeba rozumieć do dna

Ta różnica jest źródłem większości pomyłek. **θ (theta) to opór cieplny** w ścisłym sensie: $\Delta T = \theta \cdot P$ jest prawdziwe tylko wtedy, gdy przez mierzoną drogę płynie cała moc strat. Dlatego θ_{JC} mierzy się z zimną płytą wymuszającą przepływ niemal wyłącznie przez obudowę – wtedy warunek jest spełniony i θ_{JC} jest fizycznym oporem cieplnym.

Ψ (psi) to parametr korelacyjny. Mierzy się go na płycie JEDEC, gdzie ciepło rozdziela się na wiele dróg naraz. Iloczyn $\Psi \cdot P$ nie jest spadkiem na fizycznym oporze – to skalirowany współczynnik wiążący T_J z temperaturą odniesienia dla konkretnego rozdziału mocy:

$$\Psi_{JT} = \frac{T_J - T_{top}}{P}$$

Dlaczego więc Ψ działa, skoro θ_{JA} nie działa? Nota ROHM tłumaczy to wprost [7]: ponieważ SMD oddaje większość ciepła do płytki, strumień płynący ku wierzchowi obudowy jest znikomy, więc temperatura wierzchu leży blisko

temperatury złącza – różnica $T_J - T_{top}$ jest mała, a zatem i Ψ_{JT} jest małe. A skoro Ψ_{JT} jest małe, to **błąd oszacowania T_J pozostaje mały nawet wtedy, gdy Twoja płytki różni się od płytki JEDEC**. Tego θ_{JA} nigdy nie zaoferuje.

Warto dodać ważne zastrzeżenie: nawet θ_{JC} nie jest wolne od kłopotów pomiarowych – klasyczny artykuł Dutty pytał w tytule, czy opór złącze–obudowa „nie jest wciąż miernikiem” [6]. To właśnie dlatego powstała metoda przejściowa dwóch interfejsów (JESD51-14), dokładniej rozdzielająca opór wewnętrzny od oporu styku [4].

Jak wyznaczyć T_J w praktyce

Metoda A – θ_{JC} + radiator. Dla elementów mocy chłodzonych radiatorem buduje się klasyczny łańcuch oporów: złącze–obudowa, obudowa–radiator (styk) i radiator–otoczenie.

$$T_J = T_A + P \cdot (\theta_{JC} + \theta_{CS} + \theta_{SA})$$

To jedyny przypadek, w którym rachunek „z katalogu” jest wiarygodny – bo θ_{JC} jest prawdziwym oporem. Do policzenia strat i doboru elementów pomocne są firmowe narzędzia, np. TI SLUA566 [12].

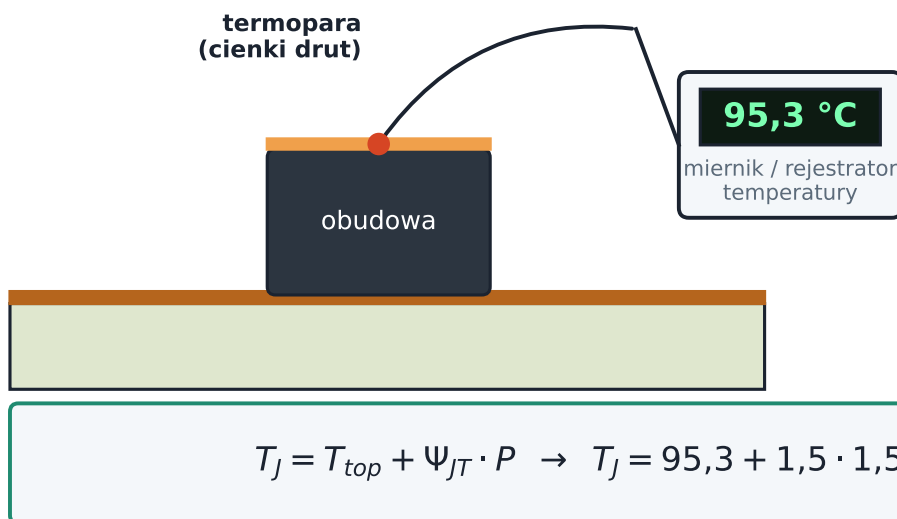
Metoda B – Ψ_{JT} i pomiar wierzchu. Najprostsza metoda „w boju”. Do wierzchu obudowy przyklej termoprzewodzącym klejem cienką termoparę (tak robi to ROHM w procedurze pomiarowej), obciąż układ realną mocą, odczytaj T_{top} i policz $T_J = T_{top} + \Psi_{JT} \cdot P$ (rysunek 6) [7]. Uwagi: cienki drut (mała masa cieplna), dobry kontakt spoiny, nie zaburzać przepływu powietrza samą termoparą, uśredniaj odczyt. Ta metoda mierzy realne warunki Twojego urządzenia – łącznie z lokalną temperaturą powietrza.

Z noty producenta – onsemi sam odsyła od θ_{JA}

W karcie sterownika bramkowego FAN3121 onsemi **nie każe liczyć T_J z θ_{JA}** . Prowadzi projektanta wzorem $T_J = P \cdot \Psi_{JB} + T_B$, opartym na temperaturze płytki. Dla $\Psi_{JB} = 42^\circ\text{C/W}$ i strat $P \approx 0,40\text{ W}$, przy odchyleniu od granicy 150°C na 80% (czyli $T_J \leq 120^\circ\text{C}$), z przekształcenia wyznacza **maksymalną dopuszczalną temperaturę płytki $T_B \approx 104^\circ\text{C}$** [9]. To sam producent, we własnej nocie, kieruje Cię od θ_{JA} ku parametrowi mierzonemu na płycie.

Praktyczny pomiar T_J metodą Ψ_{JT}

Mierzysz realną temperaturę wierzchu obudowy — poprawka Ψ jest mała



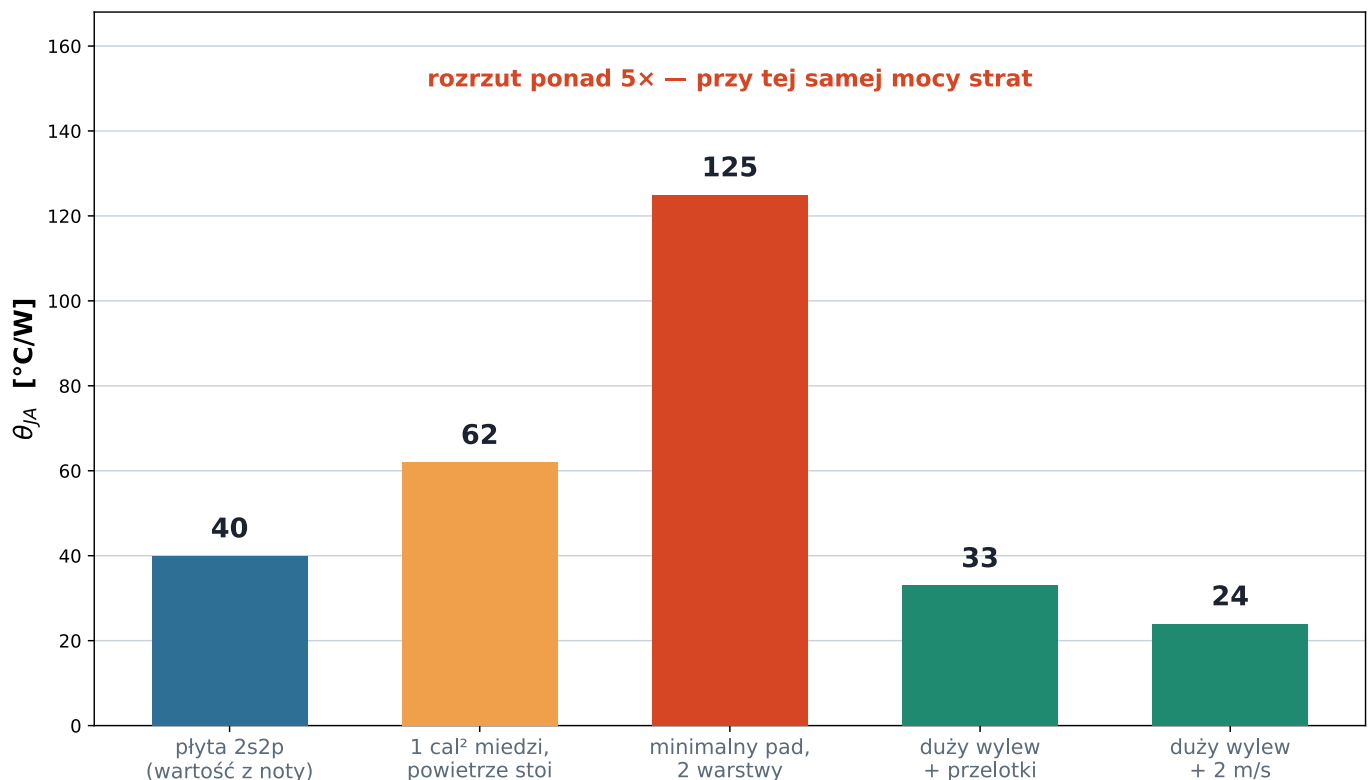
6. Pomiar T_J metodą Ψ_{JT} : cienka termopara przyklejona w środku wierzchu obudowy, realne obciążenie, mała poprawka $\Psi_{JT} \cdot P$. Wynik odzwierciedla faktyczne warunki w urządzeniu

Metoda C – symulacja. Model kompaktowy (dwuoporowy albo DELPHI) lub pełna symulacja CFD (ang. *Computational Fluid Dynamics* – obliczeniowa mechanika płynów)/FEM (ang. *Finite Element Method* – metoda elementów skończonych). Warunkiem sensownego wyniku jest wierny model Twojej płytki. Producenci dostarczają dziś dane do takich modeli – np. ROHM publikuje

raporty SPICE (ang. *Simulation Program with Integrated Circuit Emphasis*)/termiczne dla stabilizatorów [8].

Metoda D – pomiar elektryczny (TSEP, ang. *temperature-sensitive electrical parameter* – parametr elektryczny czuły na temperaturę). Najdokładniejsza, bo mierzy samo złącze. Wykorzystuje parametr elektryczny czuły na temperaturę – napięcie przewodzenia złącza p-n przy

Ta sama obudowa DPAK — pięć wartości θ_{JA}



7. Ta sama obudowa DPAK, ta sama moc strat – pięć wartości θ_{JA} różniących się ponad pięciokrotnie, zależnie od płytki i przepływu. Liczba z noty ma sens dopiero z założeniami

Reguły praktyczne • Zapamiętaj sześć zasad

- ❖ **Nie projektuj marginesu na nagłówkowym θ_{JA} .** ROHM nazywa to wprost nadużyciem – ten wzór opisuje płytkę JEDEC, nie Twoją.
- ❖ **θ_{JA} bez przypisu (jaka płyta, jaki przepływ) nie znaczy nic.** Najpierw znajdź przypis.
- ❖ **Element z radiatorem – liczb θ_{JC} + θ_{CS} + θ_{SA} .** To jedyny wiarygodny rachunek „z katalogu”.
- ❖ **SMD bez radiatora – projektuj miedź.** Pad, wylew, przelotki i liczba warstw to Twoje dzwignie.
- ❖ **Weryfikuj TJ:** Ψ_{JT} + termopara (albo Ψ_{JB} + temperatura płytki, jak zaleca onsemi), symulacja na realnej płytce albo pomiar elektryczny (TSEP).
- ❖ **θ_{JA} używaj tylko do rankingu obudów** i względnego ΔT_J , przy identycznych założeniach dotyczących płytki i otoczenia.

ustalonym prądzie pomiarowym, VBE (napięcie baza–emiter) albo RDS(on) (rezystancja dren–źródło w stanie włączenia) – skalibrowany współczynnikiem K. Dokładnie tak JEDEC (i ROHM w swojej nocie) mierzy TJ; tę samą zasadę możesz zastosować u siebie [7].

Przykład rachunkowy – jak bardzo można się pomylić

Weźmy stabilizator LDO (ang. Low-Dropout – o małym spadku napięcia) w obudowie DPAK, rozpraszający $P = 1,5$ W, przy lokalnej temperaturze wewnątrz urządzenia $T_A = 40^\circ\text{C}$. Prześledźmy cztery scenariusze (rysunek 7).

1. **Naiwnie, z nagłówka noty.** $\theta_{JA} = 40^\circ\text{C/W}$ (płyta 2s2p). $T_J = 40 + 1,5 \cdot 40 = 100^\circ\text{C}$. Margines do granicy 150°C wygląda komfortowo.
2. **Rzeczywista płytka: dwie warstwy, minimalny pad.** $\theta_{JA} \approx 125^\circ\text{C/W}$. $T_J = 40 + 1,5 \cdot 125 = 227^\circ\text{C}$. Element wchodzi w zabezpieczenie termiczne albo ginie. Te „komfortowe” 100°C były fikcją – pomyłka o 127°C .
3. **Ten sam układ z dużym wylewem i przelotkami.** $\theta_{JA} \approx 33^\circ\text{C/W}$. $T_J = 40 + 1,5 \cdot 33 \approx 90^\circ\text{C}$. Ta sama obudowa, ta sama moc – bezpiecznie. Różnicę zrobił projekt miedzi.
4. **Weryfikacja pomiarem.** Na płytce z wylewem, w rzeczywistej obudowie urządzenia, termopara

Sylwetki przywołanych ekspertów

Clemens J. M. Lasance (zm. 2025) – Philips Research, Eindhoven. Fizyk z wykształcenia (Eindhoven), od 1969 r. w Philips Research, gdzie od 1984 r. zajmował się termiką układów elektronicznych i doszedł do stanowiska principal scientist (na emeryturze od 2009 r.). Był jednym z pionierów stosowania obliczeniowej mechaniki płynów (CFD) w chłodzeniu elektroniki oraz głównym autorem monografii o kompaktowych modelach termicznych. Kierował europejskimi projektami DELPHI i PROFIT, które ugruntowały ideę modeli termicznych niezależnych od warunków brzegowych (BCI) – wprost adresującą ułomność pojedynczego θ_{JA} . Od 1995 r. był redaktorem magazynu *Electronics Cooling*; otrzymał m.in. nagrodę Harveya Rostena (2006) oraz tytuł *Significant Contributor SEMI-THERM* (2001). Jego prace o „operacyjnym” oporze cieplnym obudów są jednym z fundamentów tego artykułu.

Darwin R. Edwards – TI Fellow, Texas Instruments (obecnie Edwards' Enterprise Consulting). Ukończył fizykę na Arizona State University (1980) i niemal od razu dołączył do Texas Instruments, gdzie zbudował kompetencje firmy w charakteryzacji i zarządzaniu termiką oraz jej laboratoria termiczne. W 1999 r. został TI Fellow, a w latach 1997–2012 kierował grupą Advanced Package Modeling and Characterization. Jest współautorem pięciu norm JEDEC – w tym specyfikacji płytek testowych o niskiej i wysokiej przewodności cieplnej (JESD51-3 i -7), na których do dziś mierzy się θ_{JA} . Wraz z Hiepem Nguyenem napisał kanoniczną notę *Semiconductor and IC Package Thermal Metrics* (SPRA953), w której wprost wskazuje, że θ_{JA} jest najczęściej raportowanym i zarazem najczęściej nadużywanym parametrem cieplnym. Po odejściu z TI (2013) prowadzi firmę doradczą zajmującą się niezawodnością i termiką obudów; jest Senior Memberem IEEE (*Institute of Electrical and Electronics Engineers*) i posiada ponad 20 patentów.

na wierzchu pokazuje $T_{top} = 95,3^\circ\text{C}$ przy pełnym obciążeniu (lokalne powietrze jest cieplejsze niż nominalne 40°C). Z $\Psi_{JT} = 1,5^\circ\text{C/W}$: $T_J = 95,3 + 1,5 \cdot 1,5 \approx 97,6^\circ\text{C}$. To jest liczba, której można ufać – bo pochodzi z realnych warunków. To ta sama logika, którą stosuje onsemi, licząc T_J z Ψ_{JB} i temperatury płytki [9].

Wniosek: sama liczba θ_{JA} nie powiedziała projektantowi, którą z tych płytek ma przed sobą. Ta sama nota „uzasadniała” i bezpieczne 90°C , i śmiertelne 227°C . Dlatego – **θ_{JA} bez płytki to fikcja.**

WM, JS

Bibliografia

- [1] JESD51-1, Integrated Circuits Thermal Measurement Method – Electrical Test Method (Single Semiconductor Device).
- [2] JEDEC JESD51-2A, Integrated Circuits Thermal Test Method Environmental Conditions – Natural Convection (Still Air).
- [3] JEDEC JESD51-3/-5/-7, płytki testowe o niskiej i wysokiej efektywnej przewodności cieplnej dla obudów SMD.
- [4] JEDEC JESD51-14, Transient Dual Interface Test Method for the Measurement of Thermal Resistance Junction-to-Case.
- [5] C. J. M. Lasance, Factors Affecting the Operational Thermal Resistance of Electronic Packages, *Journal of Electronic Packaging*, vol. 122, no. 3, 2000.
- [6] V. B. Dutta, Junction-to-Case Thermal Resistance – Still a Myth?, materiały konferencyjne SEMI-THERM.
- [7] ROHM Semiconductor, θ_{JA} and Ψ_{JT} , Application Note No. 65AN114E, 2023.
- [8] ROHM Semiconductor, BD9xxN1-C SPICE Modeling Report, Application Note (stabilizatory LDO).
- [9] onsemi, FAN3121/FAN3122 – Single 9-A High-Speed Low-Side Gate Driver, karta katalogowa (FAN3121-F085/D).
- [10] onsemi, FDB3672-F085, karta katalogowa; onsemi Community – zależność θ_{JA} od projektu PCB i warunków aplikacji.
- [11] D. Edwards, H. Nguyen, *Semiconductor and IC Package Thermal Metrics*, Application Report SPRA953 (Rev. C), Texas Instruments, 2016.
- [12] Texas Instruments, Using Thermal Calculation Tools for Analog Components, Application Report SLUA566.
- [13] Texas Instruments, LMG3410 600-V 12-A Single-Channel GaN Power Stage, karta katalogowa (rewizja θ_{JA} do warunków standardowej płytki JEDEC).
- [14] OmniVision, WS4666D, karta katalogowa ($R\theta_{JA}$ oraz $R\theta_{JC}$ top/bottom).
- [15] IPC-2152/IPC-2221, wytyczne projektowe dotyczące zachowania cieplnego ścieżek i pól miedzi na PCB.
- [16] Analog Devices EngineerZone – interpretacja θ_{JA} , θ_{JC} , θ_{JB} ; ElectronicsBeliever, MOSFET Power Dissipation and Junction Temperature Calculation (materiały praktyczne).

Stabilizacja pętli trójfazowego prostownika Vienna, część 1

Trójfazowy prostownik Vienna został opatentowany przez prof. Johanna Kolar w grudniu 1993 roku, o czym opowiedział podczas wykładu plenarnego wygłoszonego na konferencji Applied Power Electronics Conference (APEC) w 2018 roku. W swojej najprostszej, jednokierunkowej realizacji prostownik wymaga trójfazowego mostka diodowego połączonego z trzema przełącznikami dwukierunkowymi, które obecnie powszechnie wykonuje się przy użyciu dwóch tranzystorów SiC połączonych antyszeregowo.

Sterowanie tymi elementami może być realizowane na różne sposoby: za pomocą pełnego sterowania histeretowego, sinusoidalnej modulacji szerokości impulsów (SPWM) lub modulacji wektora przestrzennego (SVM). W typowej aplikacji ze sterowaniem dq0 istnieją trzy pętle do stabilizacji. Jednak metody stosowane obecnie do projektowania kompensacji tych pętli mają swoje wady.

Niektóre programy korzystają z silnika automatycznego strojenia, który dostosowuje parametry kompensatora automatycznie – jest to jednak metoda prób i błędów. Co najmniej jeden program, PLECS, potrafi wygenerować małosygnałową odpowiedź całego układu przetwornika, ale wymaga to zakupu licencji. Choć projektanci mogliby przeprowadzić analizę małosygnałową układu prostownika Vienna, dla większości z nich proces ten jest zbyt skomplikowany.

Aby zaradzić tej sytuacji, opracowałem model uśredniony złożony z sześciu modeli przełącznika PWM, który dostarcza odpowiedzi prądu przemiennego potrzebnych do stabilizacji przetwornika w stosunkowo prosty i szybki sposób. W pierwszej części artykułu zbadamy strategię opartą na tym modelu uśrednionym do kompensacji pętli sterowania prostownika Vienna przy użyciu układów analogowych. W drugiej części zaktualizujemy tę metodę dla przetwornika pracującego z kompensatorami cyfrowymi.

W części 1 rozpoczniemy od przetwornika T-type, który stanowi dobry punkt wyjścia do wyjaśnienia działania prostownika Vienna przez sprowadzenie go do realizacji jednofazowej. Posiadanie modelu uśrednionego ułatwi



Christophe Basso, urodzony w 1965 roku, to uznany ekspert w dziedzinie energoelektroniki, specjalizujący się w projektowaniu zasilaczy impulsowych oraz stabilizacji pętli sterowania. Posiada tytuł Senior Member IEEE oraz jest autorem licznych patentów. Swoją edukację techniczną rozpoczął na Uniwersytecie w Montpellier we Francji, a tytułu magistra inżyniera elektrotechniki ze specjalnością w energoelektronice uzyskał na uczelni Institut National Polytechnique w Tuluzie.

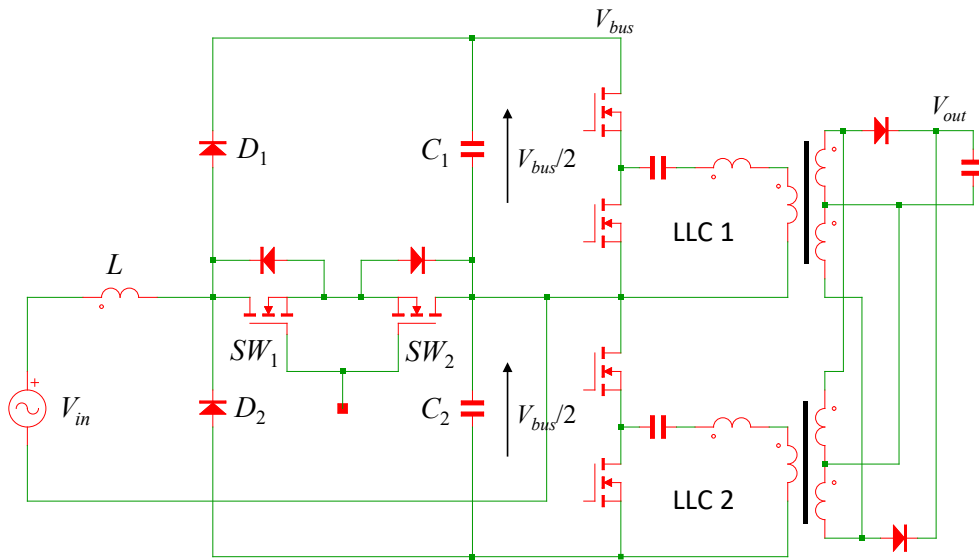
Jego bogata ścieżka zawodowa obejmuje pracę w kluczowych instytucjach i firmach technologicznych. Początkowo, przez dziesięć lat, pracował jako projektant układów zasilania w Europejskim Ośrodku Promieniowania Synchrotronowego w Grenoble. Następnie przez dwa lata był związany z działem półprzewodników firmy Motorola. Najdłuższy, trwający niemal 24 lata etap jego kariery, to praca w firmie onsemi w Tuluzie. Pełnił tam prestiżową funkcję Technical Fellow oraz dyrektora do spraw inżynierii produktu, stając się wiodącą postacią w projektowaniu i wprowadzaniu na rynek nowoczesnych kontrolerów przełączających. Po zakończeniu współpracy z onsemi kontynuował karierę w branży dystrybucji komponentów elektronicznych, gdzie wspiera wdrażanie innowacyjnych rozwiązań zasilających.

Basso jest postacią niezwykle cenioną w globalnym środowisku inżynierskim. Znany jest jako autor jedenastu książek technicznych, w tym popularnych publikacji na temat metod szybkiej analizy obwodów oraz symulacji układów zasilających. Jego wkład w branżę to nie tylko teoria, ale przede wszystkim praktyczne narzędzia. Opracował i udostępnił obszerne, darmowe biblioteki modeli symulacyjnych dla oprogramowania SPICE i LTspice, które stanowią ogromne ułatwienie dla projektantów na całym świecie. Ponadto regularnie dzieli się swoim doświadczeniem, występując jako prelegent na międzynarodowych seminariach i konferencjach branżowych.

uzyskanie odpowiedzi prostownika Vienna przy użyciu LTspice. Niniejsze opracowanie krótko omówi koncepcję transformacji Clarka, która jest podstawą sterowania dq0, a następnie pokaże jak model uśredniony z sześcioma przełącznikami PWM może być wykorzystany do kompensacji trzech pętli prostownika Vienna pracującego pod sterowaniem dq0.

Jednofazowy przetwornik boost T-type

Rysunek 1 przedstawia trójpoziomowy przetwornik boost pracujący z przełącznikiem dwukierunkowym. Konfiguracja ta nosi nazwę przetwornika T-type, a napięcie wyjściowe V_{bus} dzieli się między dwa kondensatory wyjściowe C_1 i C_2 . Dla wejściowego napięcia skutecznego 265 V przetwornik ten będzie regulował napięcie



1. Przetwornik trójpoziomowy doskonale nadaje się do kaskadowania i redukcji napięć na półprzewodnikach. Tu dwa przetworniki LLC są połączone równolegle i zasilane z dwóch różnych napięć pierwotnych. Inne typy przetworników mogą być tu również stosowane do konwersji DC-DC z izolacją

na poziomie 800 V DC, umożliwiając łączenie dwóch przetworników zasilanych przez szyny 400 V. Takie podejście pozwala projektantowi wybrać tranzystory o napięciu przebicia 650 V, podczas gdy przy klasycznym przetworniku dwupoziomowym z jedną szyną 800 V wymagane byłyby typy o $BV_{DSS} = 1200$ V.

W tej konfiguracji przełączniki SW_1 i SW_2 są połączone antyszeregowo, zapewniając dwukierunkowość. Na przykład, gdy V_{in} jest dodatnie, przełączniki SW_1/SW_2 są spolaryzowane i napięcie wejściowe pojawia się na indukcyjności L . Prąd narasta ze zboczem:

$$S_{on} \approx \frac{V_{in}}{L}$$

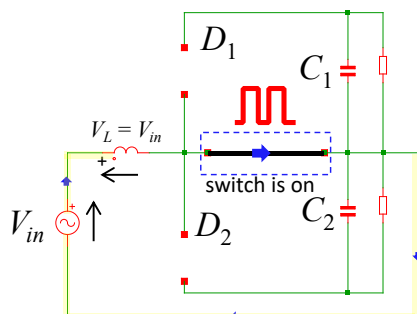
Gdy przełączniki otwierają się, prąd indukcyjny dalej płynie w tym samym kierunku, lecz teraz swobodnie

płynie przez diodę D_1 . Napięcie na indukcyjności odwraca się, redukując prąd ze zboczem:

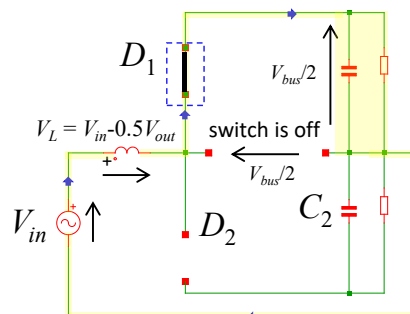
$$S_{off} \approx \frac{V_{in} - 0,5 \cdot V_{out}}{L}$$

Prąd krąży w kondensatorze C_1 i obciążeniu. Podczas tego zdarzenia dwa przełączniki blokują $V_{bus}/2$, czyli 400 V przy napięciu szyny 800 V. Gdy polaryzacja napięcia wejściowego zmienia się, sekwencja powtarza się: dioda D_2 uczestniczy w swobodnym obiegu, a prąd zasila kondensator C_2 i obciążenie.

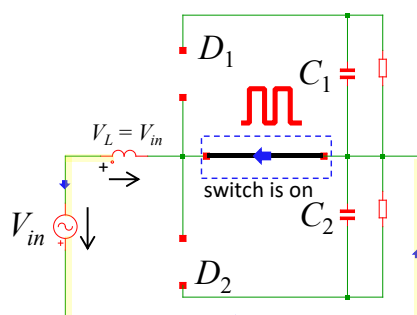
Diody D_1 i D_2 mogą być zastąpione dwoma synchronicznymi przełącznikami dla poprawy sprawności. W takim przypadku konieczne są dodatkowe układy logiczne do dekodowania polaryzacji linii wejściowej i odpowiedniego kierowania wzorców przełączania



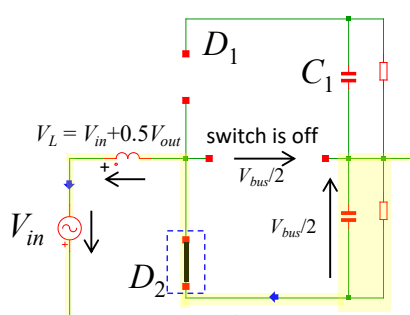
Positive ½ cycle – magnetization phase



Positive ½ cycle – demagnetization phase

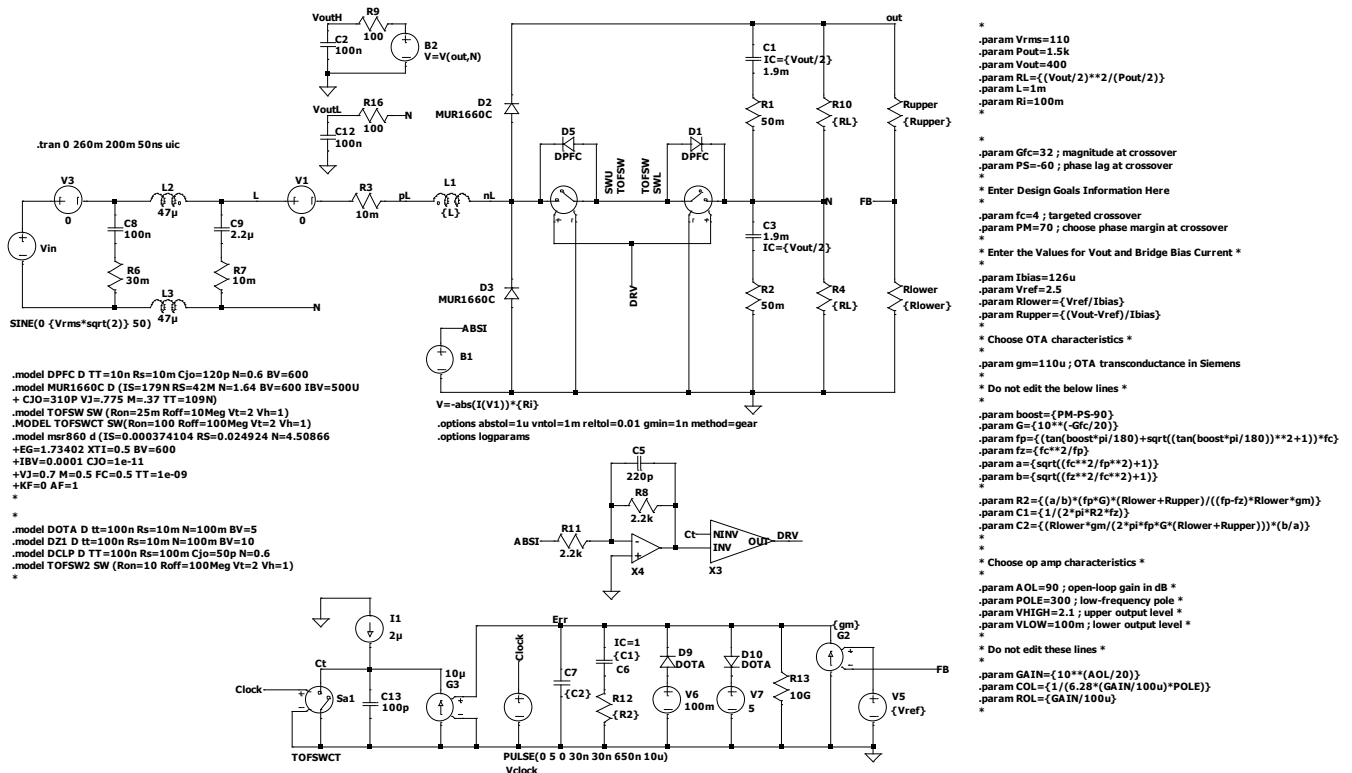


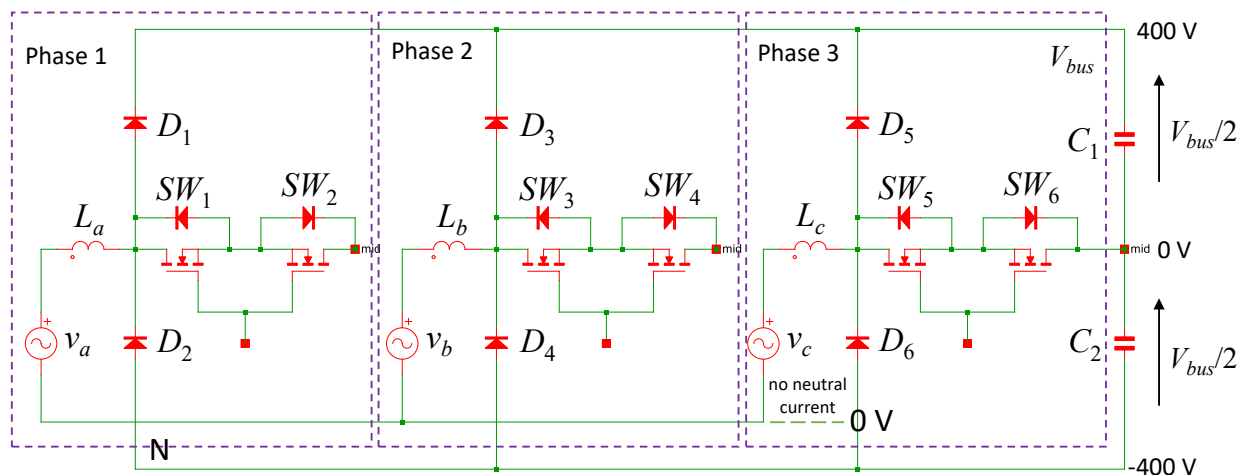
Negative ½ cycle – magnetization phase



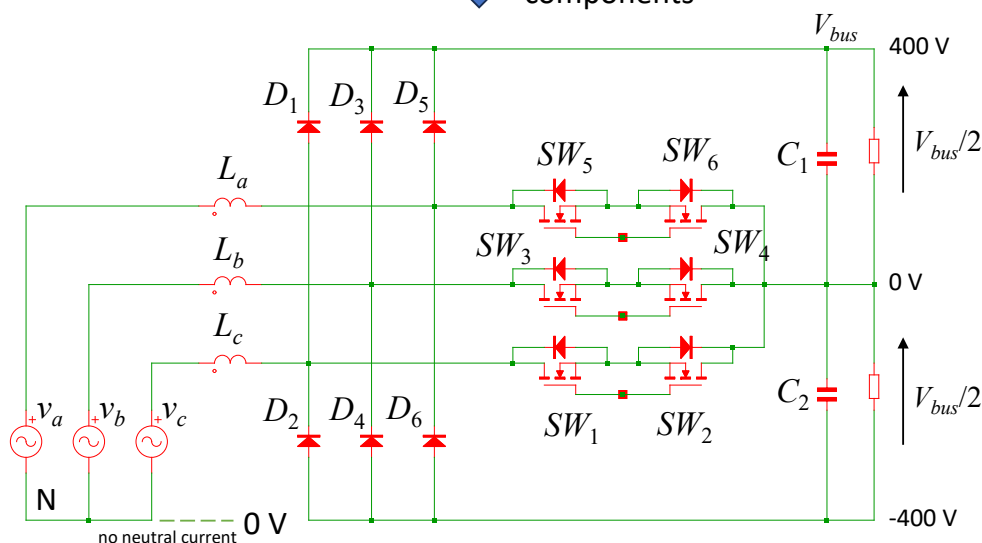
Negative ½ cycle – demagnetization phase

2. Przełącznik dwukierunkowy przewodzi w obu kierunkach w zależności od polaryzacji linii wejściowej

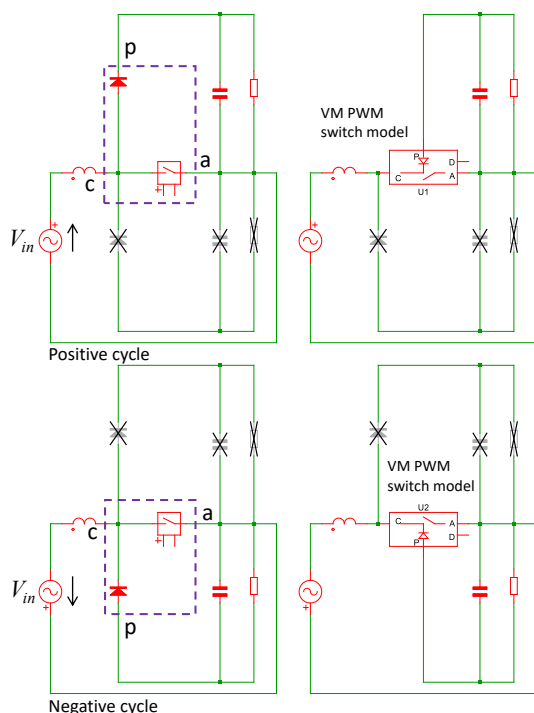




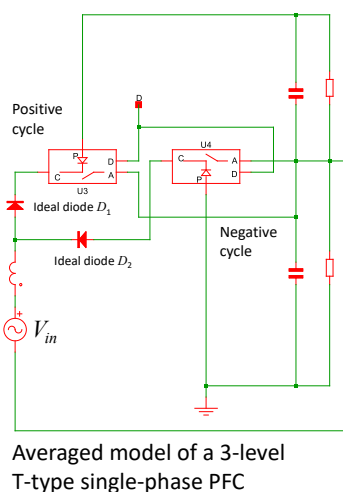
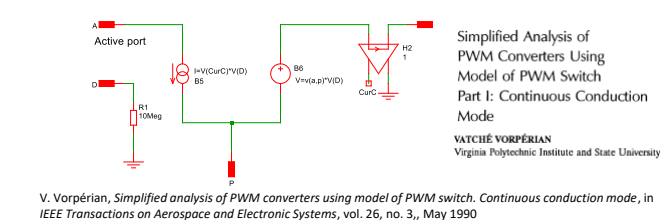
↓ Rearrange the components



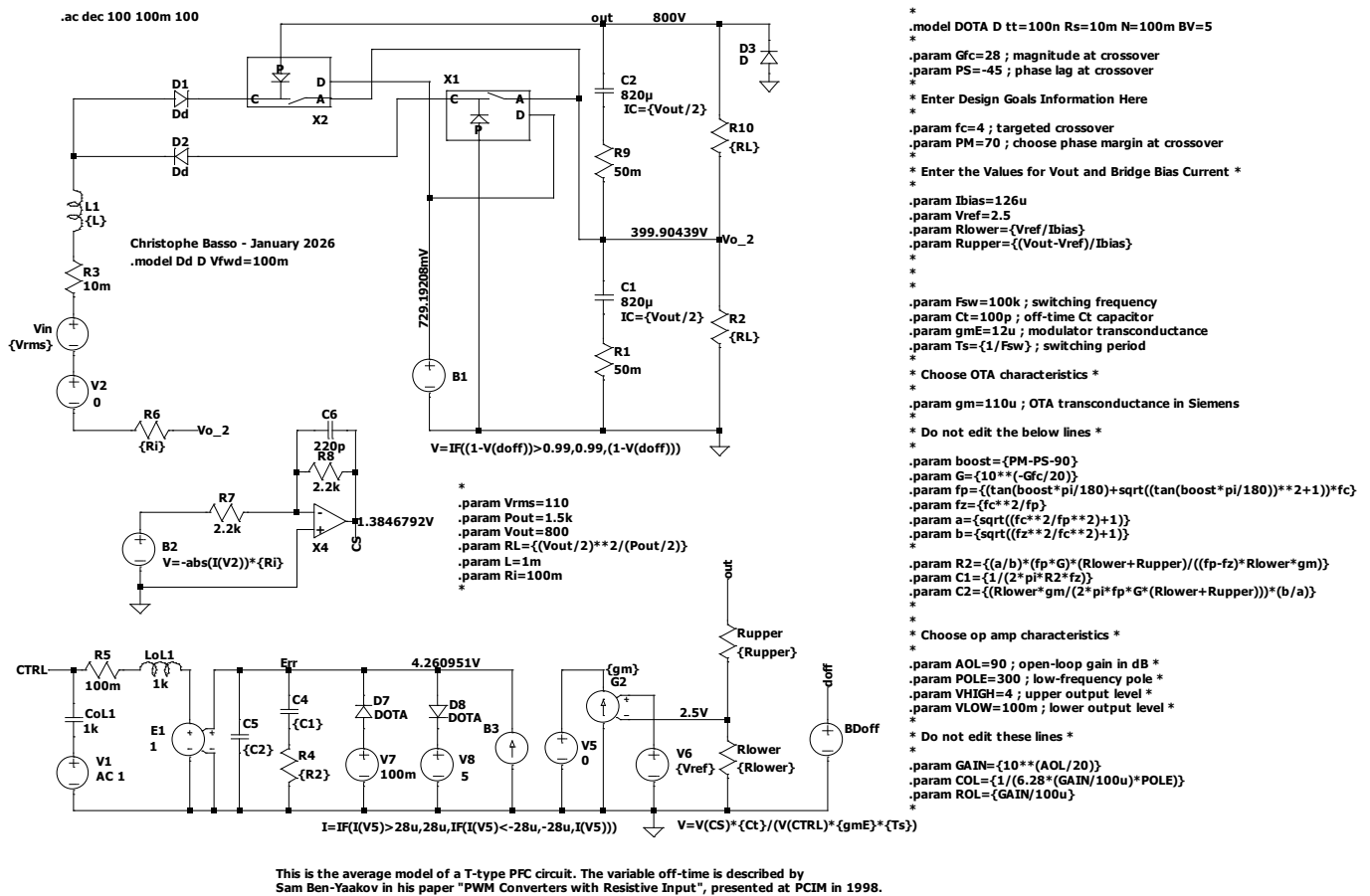
5. Zebranie trzech przetworników T-type i przeorganizowanie symboli elementów prowadzi do konfiguracji prostownika Vienna



Combine all models



6. Najpierw identyfikujemy komórkę komutacyjną podczas dodatniego półokresu, a następnie drugą podczas ujemnej części V_{in} . Składając dwa modele przetwornika PWM z szeregową diodą, uzyskujemy model uśredniony przetwornika

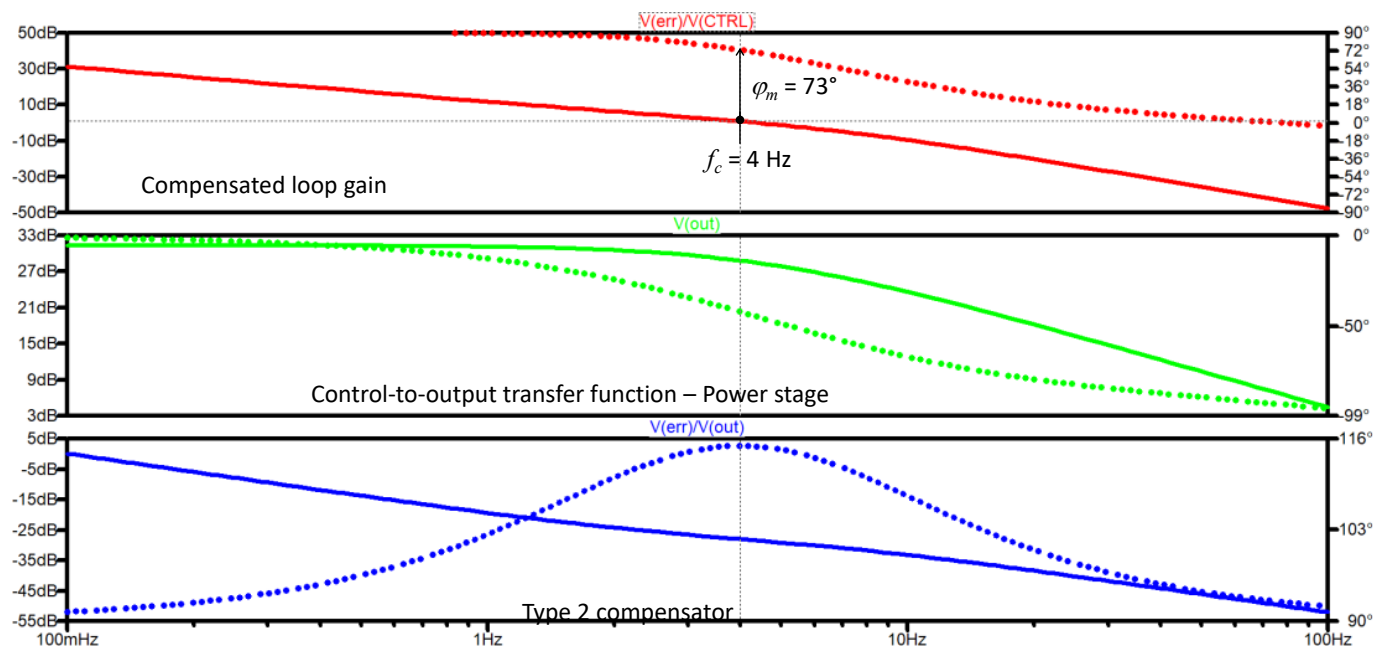


7. Modele uśrednione symulują szybko w analizie AC. Jest to falownik 800 V i punkty polaryzacji są prawidłowe

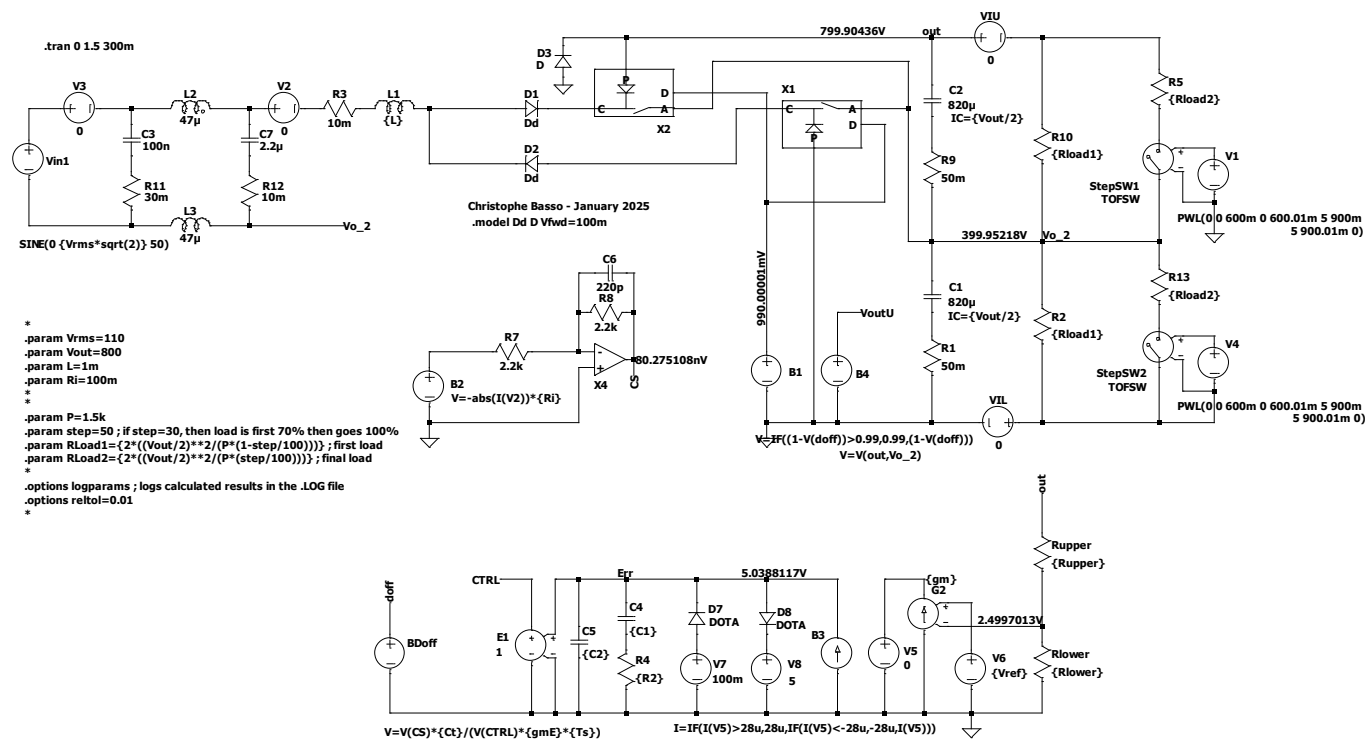
uzyskania ujemnego wyjścia przy użyciu funkcji wartości bezwzględnej (ABS). Źródło analogowe B₁ realizuje tę operację i skaluje prąd współczynnikiem rezystancyjnym R_i. Wynikowe napięcie wyjściowe przechodzi przez wysokopasmową pętlę prądową zbudowaną wokół X₄, której wyjście steruje komparatorem X₃ dla pracy PWM. Prąd ładowania kondensatora określającego timing jest sterowany

przez kompensator monitorujący napięcie błęd i dostosowujący poziom mocy dostarczanej przez przetwornik.

Ten przykład pracuje przy napięciu wejściowym 110 V i dostarcza napięcia szyny 400 V, równo podzielonego między kondensatory C₁ i C₂ przy symetrycznie obciążonych szynach. Symulacja przez 200 ms trwała 100 s na moim starym komputerze.



8. Ta symulacja AC potwierdza częstotliwość przecięcia 4 Hz przy wybranym punkcie pracy



This is the average model of a T-type PFC circuit. The variable off-time control is described by Sam Ben-Yaakov in his paper "PWM Converters with Resistive Input", presented at PCIM in 1998.

9. Ten model uśredniony pozwala na przeprowadzenie analizy przejściowej z sinusoidalnym źródłem

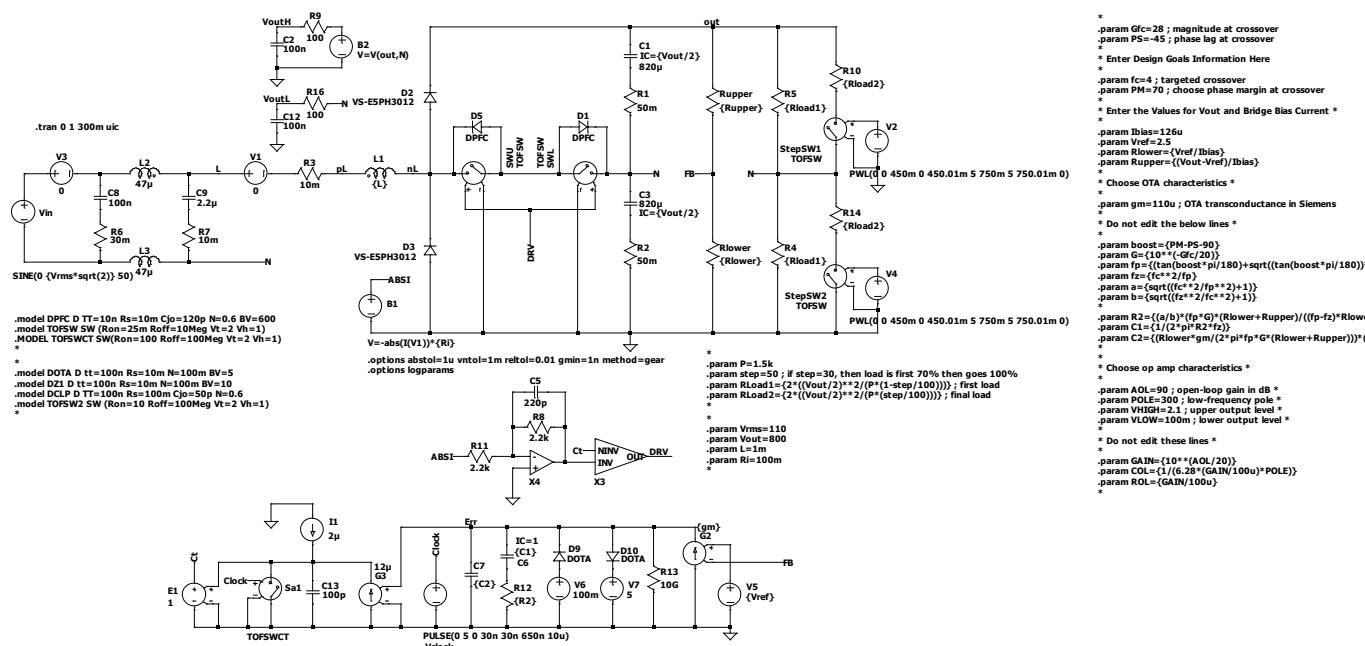
Od jednofazowego T-type do trójfazowego Vienna

Powyższa analiza jest istotna, ponieważ prostownik Vienna można rozłożyć na trzy przetworniki T-type, po jednym na każdą fazę. Rysunek 5 ilustruje, jak połączenie tych indywidualnych układów T-type prowadzi do konfiguracji Vienna. W strukturze Vienna dwa przełączniki SW₁ i SW₂ (dla fazy a) są sterowane jednocześnie, co znacznie upraszcza sterowanie. W przykładach

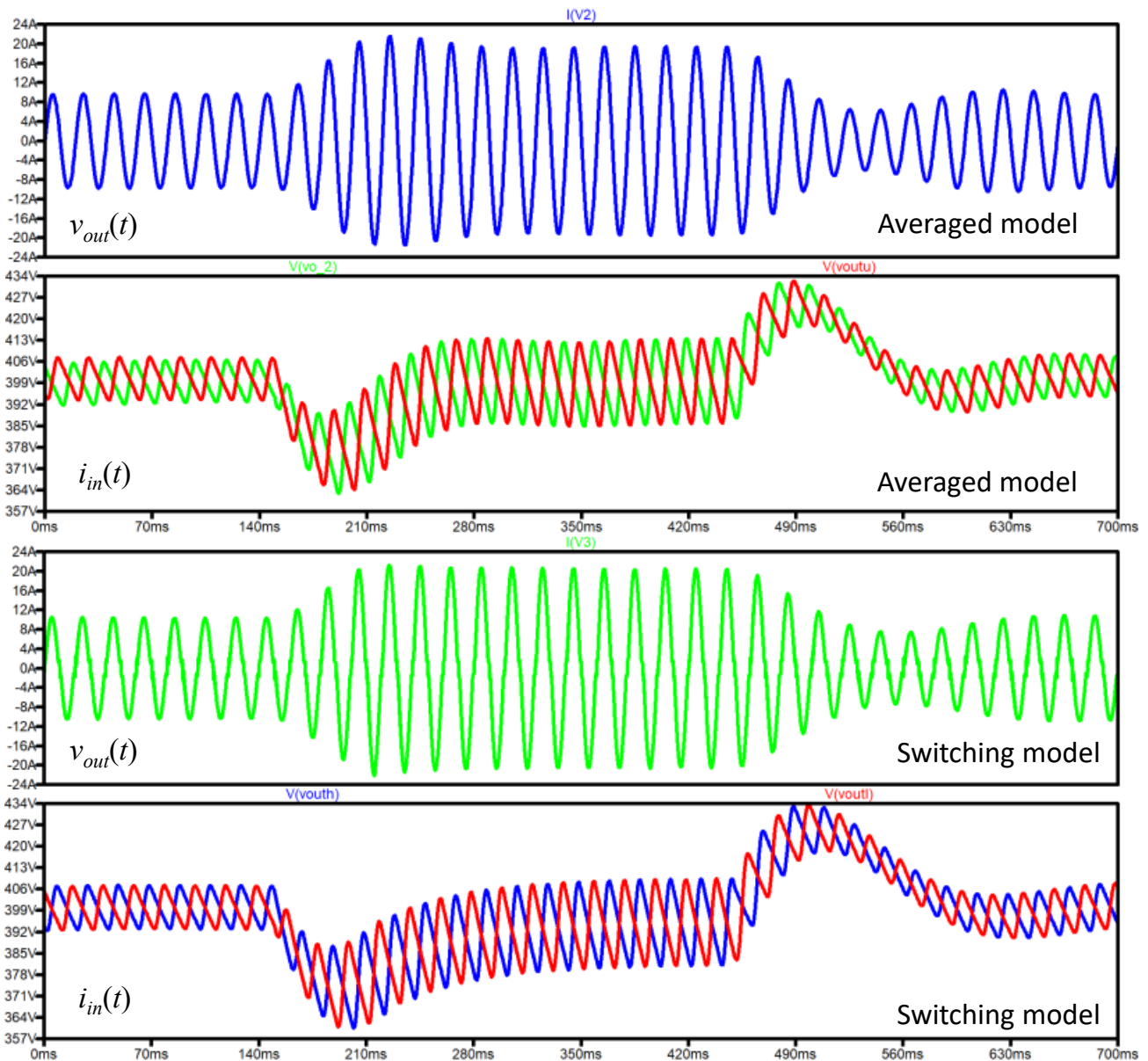
będę używał SPWM, którego implementacja to centrowana trójfazowa komutacja, identyczna z SVM typ 2.

Uśredniona symulacja jednofazowego PFC T-type

Analiza w przestrzeni stanów prostownika Vienna jest złożona. Zamiast tego skorzystam z modelu przełącznika PWM wprowadzonego przez Vatché Vorpériana w 1986 roku. Podczas dodatniego półokresu



10. Można porównać odpowiedź przejściową modelu uśrednionego z modelem przetwarzającym cykl-po-cyklu i sprawdzić, czy wyniki są wystarczająco zbliżone



11. Można porównać odpowiedź modelu uśrednionego z odpowiedzią wersji cykl-po-cyku i sprawdzić, czy wyniki są wystarczająco zbliżone

źródła wejściowego V_{in} aktywne są dioda D_1 i przełączniki SW_1/SW_2 , które podwyższają napięcie. W tym czasie dioda D_2 jest zablokowana. Podczas ujemnego półokresu aktywne są dioda D_2 i przełączniki SW_1/SW_2 , podczas gdy D_1 milczy. Zbierając te podukłady, możemy zauważyć, że dla pojedynczej komórki T-type potrzebne są dwa modele przełącznika PWM.

Do analizy transmitancji sterowanie–wyjście PFC zwykle przyjmujemy napięcie wejściowe DC równe skutecznej wartości linii AC, przy której chcemy przeprowadzić analizę. **Rysunek 7** przedstawia typowy układ stosowany do kompensacji T-type PFC. Po kilku sekundach uzyskujemy wyniki AC widoczne na **rysunku 8**. Kompensator typu 2 jest automatycznie projektowany dla osiągnięcia punktu przecięcia i zapasu fazy: 4 Hz i 70° w tym przypadku.

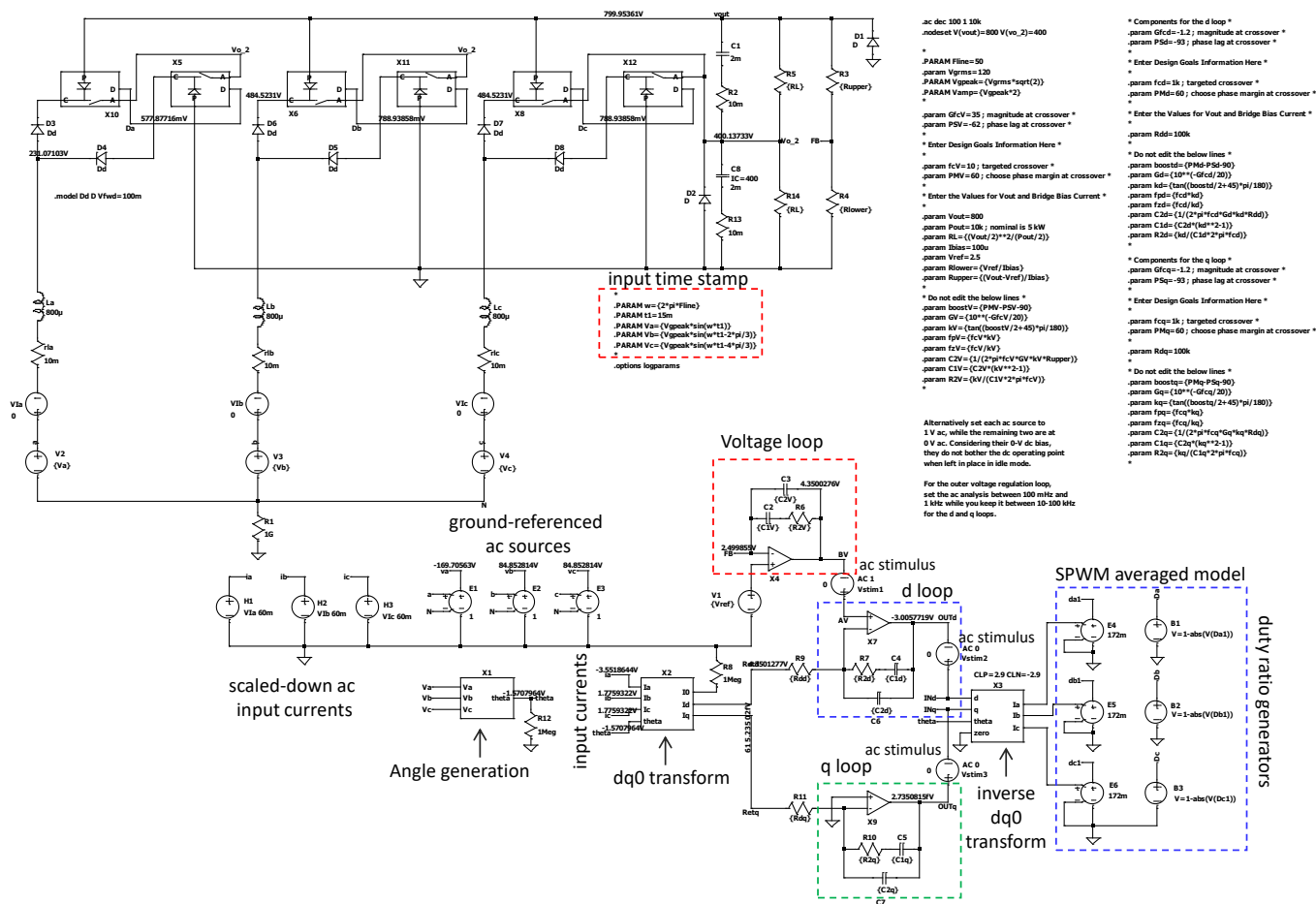
Mogę teraz uruchomić schodkowy przebieg przejściowy na modelu uśrednionym i porównać jego odpowiedź z wersją cykl-po-cyku, przy identycznie skompensowanych

przetwornikach. W tym przykładzie wyjście 800 V jest schodkowane od 50% do 100% nominalnej mocy 1,5 kW. Obie symulacje potwierdzają bardzo stabilną odpowiedź.

Uśredniona symulacja AC trójfazowego prostownika Vienna

Model uśredniony prostownika Vienna zawiera teraz sześć modeli przełącznika PWM i pełny schemat modulacji SPWM oparty na sterowaniu $dq0$. Obliczenie punktu pracy DC może stanowić wyzwanie dla silnika symulacyjnego. Testowałem ten układ zarówno ze wszystkimi sześcioma modelami przełącznika PWM, jak i z wariantem zawierającym tylko trzy z nich.

Dlaczego możemy zredukować złożoność? Ponieważ mamy trzy fazy, których chwilowe napięcia są przesunięte o 120°. W związku z tym, zależnie od wybranego punktu pracy linii wejściowej, napięcia fazowe mogą przyjmować różne polaryzacje, aktywując lub dezaktywując model poprzez



12. Uśredniony model prostownika Vienna jest dość skomplikowany, gdyż zawiera teraz sześć modeli przetwornika PWM

skojarzoną z nim diodę. Modele szeregowo z zablokowaną diodą można tymczasowo wykluczyć.

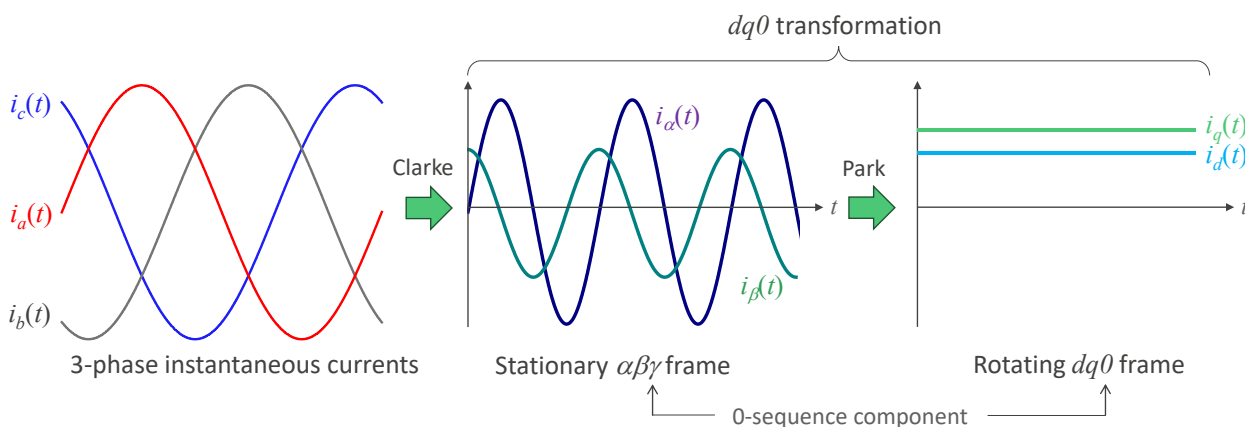
Jeśli łatwo jest wybrać wartość DC dla analizy AC w układzie jednofazowym, jak dobieramy napięcia w przypadku trzech faz? Sztuczka polega na przyjęciu znacznika czasu t_1 i wyznaczeniu chwilowej wartości każdej fazy w tym punkcie. Przetestowałem $t_1 = 15$ ms, co daje $V_a = -170$ V i $V_b = V_c = 84,85$ V.

Sterowanie mocą czynną i bierną

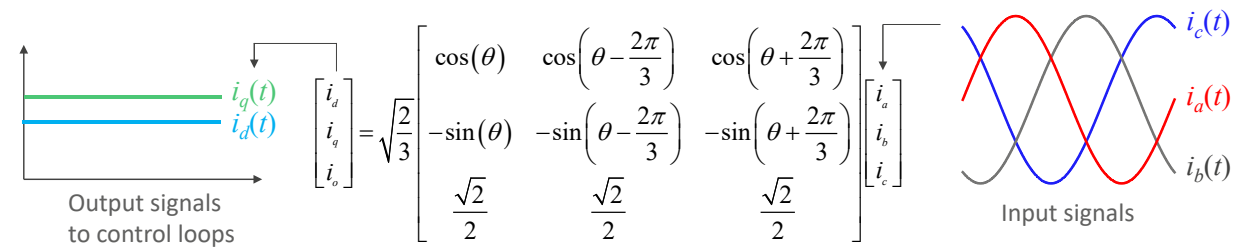
Jak opisałem we wstępie, w tym PFC istnieją trzy aktywne pętle. Istotną różnicą od jednofazowego PFC jest brak

tętnień napięciowych 100 Hz lub 120 Hz na kondensatorze wyjściowym. Dzięki temu można rozszerzyć częstotliwość przecięcia f_c poza kilka herców dopuszczalnych w aplikacjach jednofazowych. Typowe wartości dla trójfazowego PFC to $f_c = 80-100$ Hz.

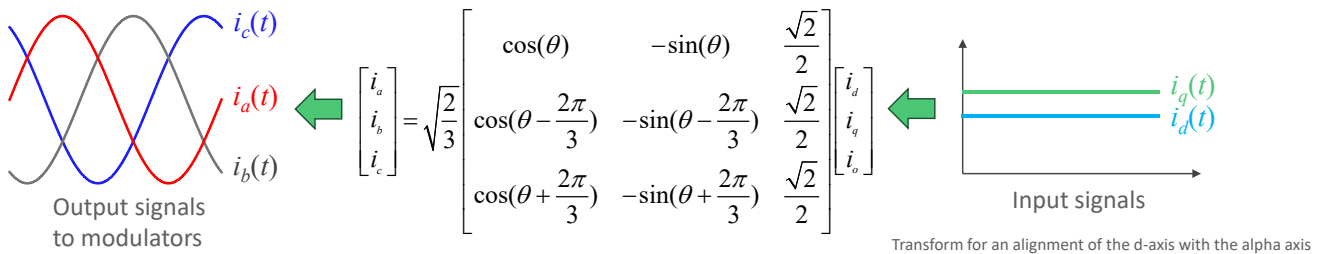
W jednofazowym PFC wyjście wzmacniacza błędów steruje bezpośrednio nastawnikiem prądu indukcyjności lub zmienia współczynnik wypełnienia. W tym trójfazowym układzie wzmacniacz błędów będzie sterować mocą czynną przez wejście pętli d, podczas gdy inna pętla z trzecim wzmacniaczem błędów monitoruje moc bierną q i utrzymuje ją na poziomie zera.



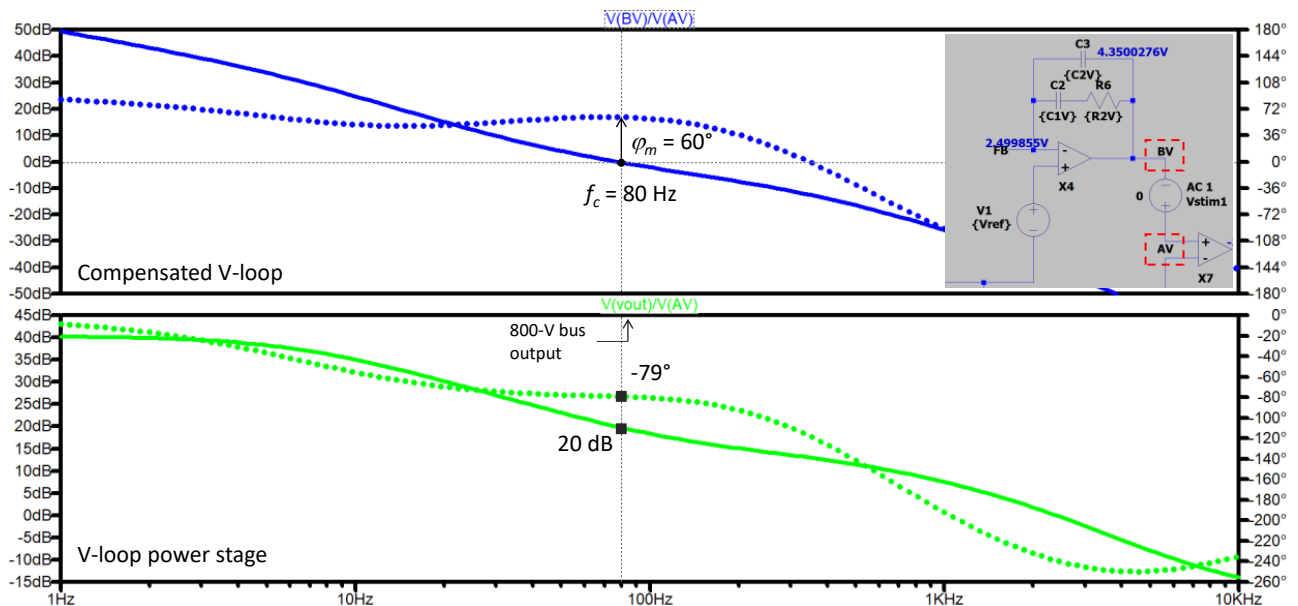
13. Sterowanie mocą z trzema sinusoidalnymi źródłami komplikuje układ



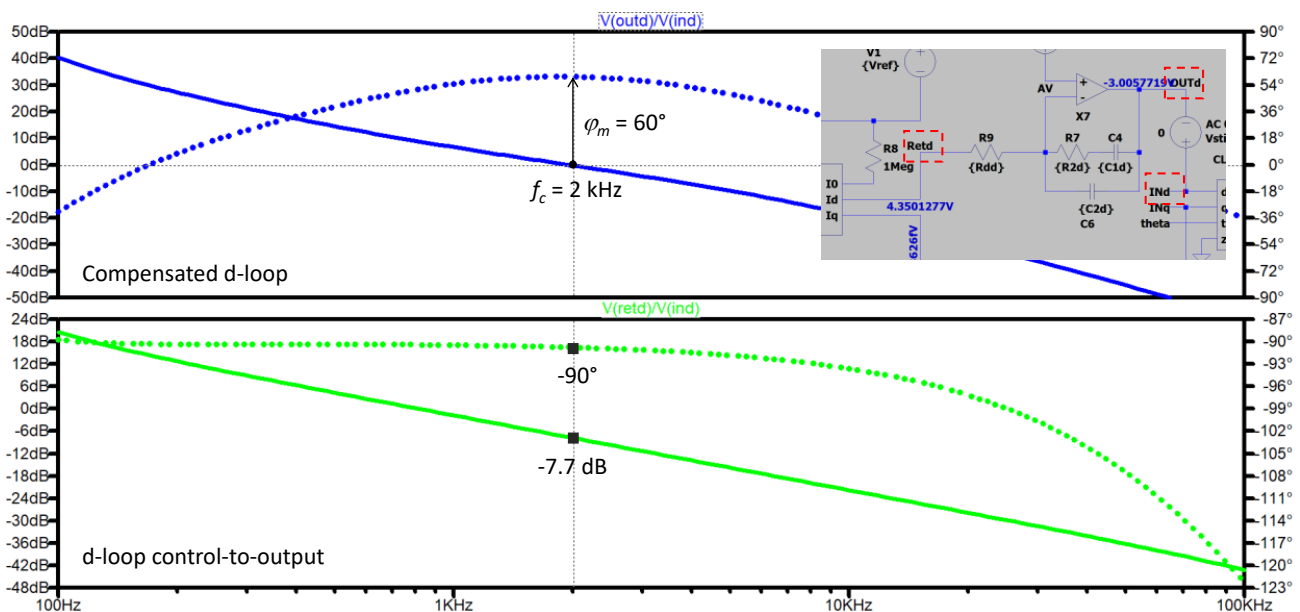
✓ The inverse $dq0$ reconstructs the corrected waveforms from regulated d and q values



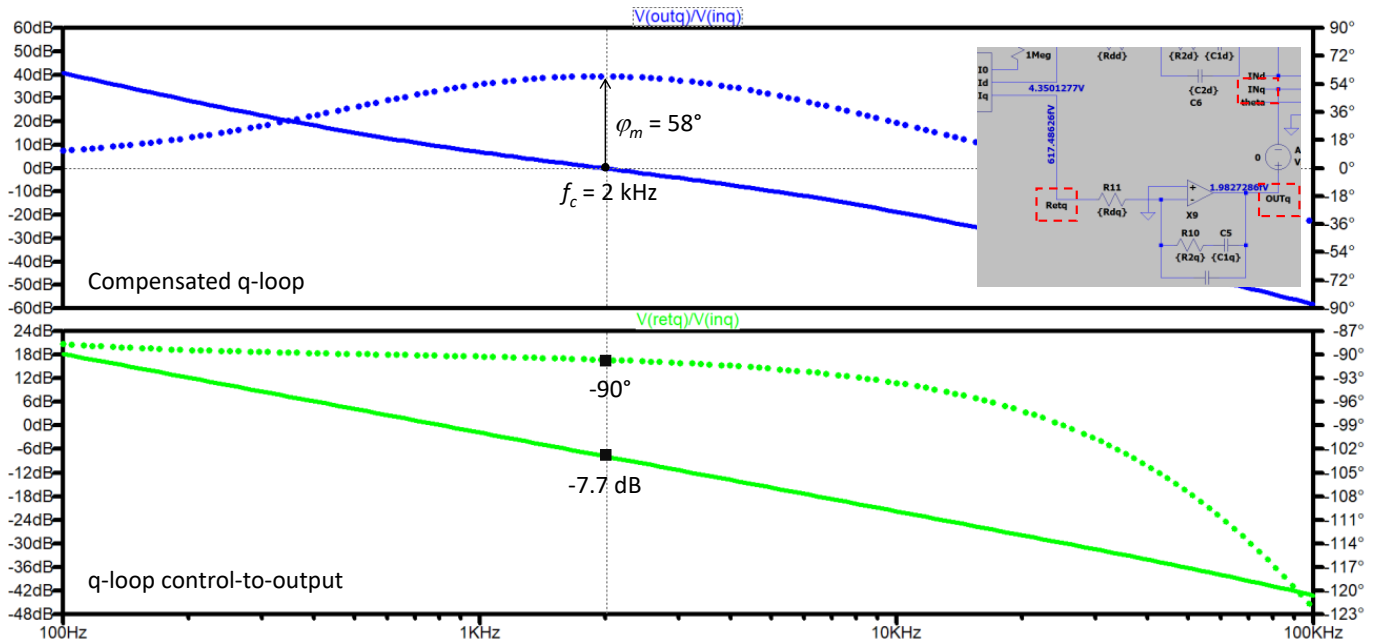
14. Transformacja dq0 jest niezmienniczym matematycznym narzędziem mocy



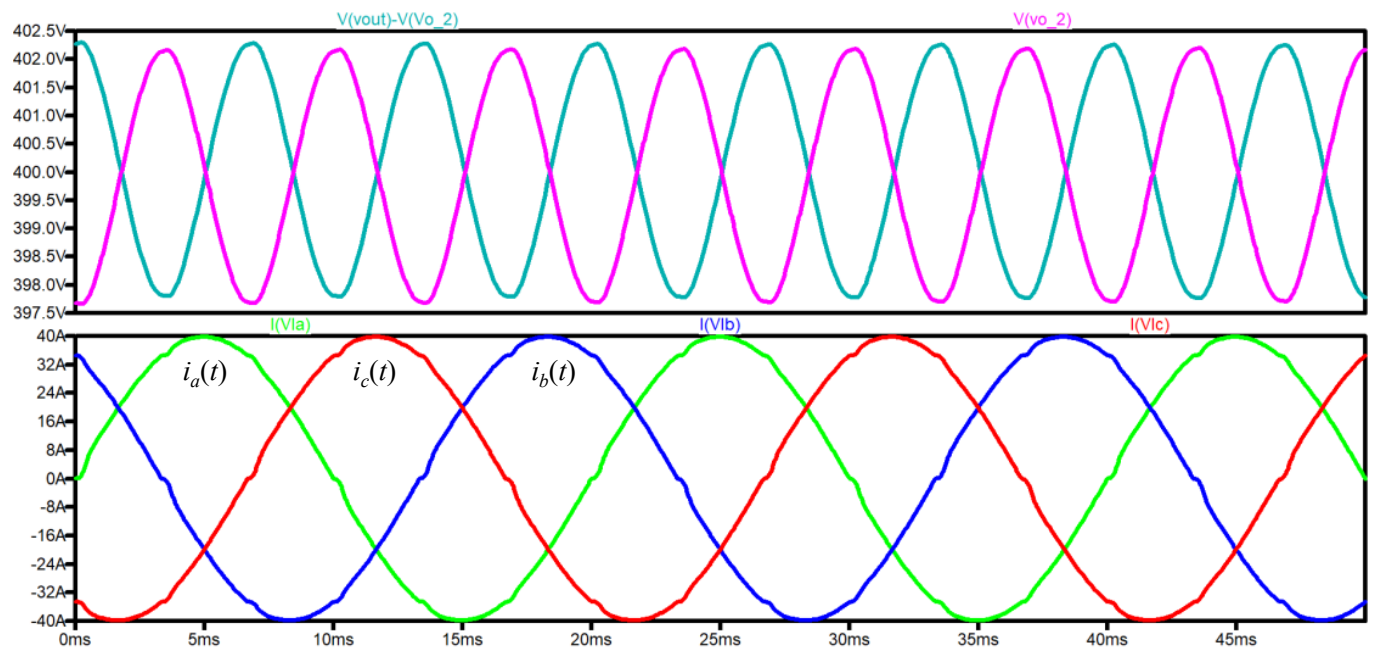
15. Pętla napięciowa jest kompensowana dla częstotliwości przecięcia 80 Hz



16. Pętla d jest kompensowana dla częstotliwości przecięcia 2 kHz



17. Pętla q jest kompensowana dla częstotliwości przecięcia 2 kHz



18. Prądy wejściowe są dobrze sinusoidalne, a obie szyny pozostają stabilne przy 400 V

Rysunek 13 przedstawia trzy sinusoidalne przebiegi prądowe dostarczane przez sieć. Alternatywne podejście polega na wektorowym sumowaniu trzech prądów wejściowych w celu syntezy jednego wektora. Transformacja matematyczna z trójwymiarowego obracającego się układu abc do dwuwymiarowego stacjonarnego układu nosi nazwę transformacji Clarka. Stosując następnie transformację Parka, wyodrębniamy składowe d i q prądów wejściowych jako wartości ciągłe – i_d i i_q – które mogą być przetwarzane przez klasyczną pętlę analogową lub cyfrową. Kąt θ jest uzyskiwany ze źródeł wejściowych za pośrednictwem pętli SRF-PLL i synchronizuje proces dq0.

W klasycznym PFC można trwale ustawić q na zero lub sterować jego wartością dla celów tworzenia sieci. Moc

czynna P, bierna Q i pozorna S wyrażone przez składowe dq0 przyjmują postać:

$$P = v_d \cdot i_d + v_q \cdot i_q + v_0 \cdot i_0 \text{ [W]}$$

$$Q = v_d \cdot i_d - v_q \cdot i_q \text{ [VAR]}$$

$$S = P + jQ \text{ [VA]}$$

Odpowiedź AC prostownika Vienna

W układzie pokazanym na rysunku 12 widoczne są trzy źródła AC szeregowo z każdym wzmacniaczem błęd. Jeśli chcemy testować pętlę napięciową, włączamy bodziec

AC V_{stim1} (AC 1), podczas gdy dwa pozostałe są wyłączone (AC 0). Dla pętli d, V_{stim2} jest włączone jako jedyne, a dla pętli q, tylko V_{stim3} jest aktywnym bodźcem.

Pętla napięciowa jest kompensowana dla częstotliwości przecięcia 80 Hz. Pętla d pracuje na identycznej zasadzie, lecz z przecięciem 2 kHz. Pętla q jest kompensowana również na 2 kHz, co jest łatwe, ponieważ odpowiedzi pętli d i q są identyczne. Prostownik Vienna jest teraz całkowicie skompensowany.

Wnioski

Prostownik Vienna był tematem wielu publikacji, jednak nie udało mi się znaleźć artykułu technicznego szczegółowo opisującego prostą metodę kompensacji jego pętli. Poprzez reprezentację trójfazowego prostownika Vienna za pomocą

sześciu modeli przełącznika PWM możliwe jest wyznaczenie odpowiedzi AC trzech pętli w sposób prosty i szybki.

Wszystkie te przykłady symulacji są dostępne do pobrania z mojej strony internetowej. Plik ZIP zawiera ponad 200 gotowych szablonów obejmujących wiele różnych przetworników impulsowych, w tym jednofazowe i trójfazowe układy PFC.

Artykuł opublikowany za uprzejmą zgodą autora Christophe'a Basso i wydawcy magazynu „How2Power” Davida Morrisona.

Tłumaczenie artykułu: Christophe Basso, „Future Electronics”, Tuluza, Francja „How2Power Today”, marzec 2026. Wszelkie prawa zastrzeżone. Etykiety i opisy na rysunkach pozostawiono w oryginalnej wersji angielskiej.

Literatura

1. „Vienna Rectifier and Beyond” – wykład plenarny J.W. Kolar, APEC 2018.
2. „PWM Converters with resistive input” – S. Ben-Yaakov, I. Zehser, IEEE Trans. Ind. Electronics, t. 45, 1998.
3. „Intuitive Explanation of the Three-Phase Vienna Rectifier” – S. Ben-Yaakov, YouTube, 2022.
4. „Three-Phase PFC Rectifier and AC-AC Converter Systems” – J. Kolar i in., APEC 2011.
5. „Vienna Rectifier & Beyond” – J. Kolar i in., APEC 2018.
6. „Modeling and Control of Three-Phase PWM Converters” – D. Boroyevic, PCon 2008.
7. „Simplified analysis of PWM converters using model of PWM switch” – V. Vorpérian, IEEE Trans. AES, t. 26, nr 3, 1990.
8. „Deriving the Control-to-Output Transfer Function of the Weinberg Converter” – C. Basso, How2Power Today, październik 2024.
9. „Novel Reduced-Order Small-Signal Model of a Three-Phase PWM Rectifier” – D. Boroyevic i in., IEEE Trans. Power Electronics, t. 13, nr 3, 1998.
10. Edith Clarke, Wikipedia.
11. „Introduction to 3-Phase PFC – Part I i II” – C. Basso, POWERSIMTOF, luty 2024.
12. POWERSIMTOF – strona internetowa Christophe'a Basso.
13. Plik ZIP z szablonami przetworników.

REKLAMA

m.technik
Ciekawi świata są zawsze młodzi

przejrzyj i kupisz na stronie
www.ulubionykiosk.pl

CHERCHER LA FEMME: Rosalind Franklin
Kobieta, która zobaczyła DNA – i której nikt nie wierzył

lipiec-sierpień • 07-08 2026
www.milodytechnik.pl

m.technik
Ciekawi świata są zawsze młodzi

SINGULARITY
KONIEC CYWILIZACJI, JAKĄ ZNAMY

Science fiction w „Młodym Techniku”
Mariusz Kocharński: Echo z Jebel Irhoud

19,90 zł

Nowoczesna technika pomiarowa wspiera precyzyjną diagnostykę i serwis

W dobie transformacji technologicznej technika pomiarowa staje się dla przedsiębiorstw coraz ważniejszym czynnikiem strategicznym – zarówno w obszarze rozwoju, jak i produkcji. Precyzyjne gromadzenie danych pomiarowych, ich dokumentowanie oraz integracja z procesami cyfrowymi pozwalają ograniczać ryzyko przestojów, podnosić jakość produktów i zmniejszać zużycie energii. Conrad Sourcing Platform odpowiada na te zmiany dzięki szerokiemu asortymentowi oraz dopasowanym usługom i modułom platformowym. W rozmowie Harald Lehner, ekspert ds. techniki pomiarowej w Conrad Electronic, omawia najważniejsze trendy rynkowe, oczekiwania klientów oraz nowoczesne rozwiązania wspierające pomiary i analizę danych.

Jak ocenia Pan obecne znaczenie rynku techniki pomiarowej, zwłaszcza dla przemysłu?

Harald Lehner: Obowiązki związane z dokumentacją w przedsiębiorstwach przemysłowych stają się coraz bardziej szczegółowe i wymagające. Jako partner zakupowy wspieramy klientów w skutecznym realizowaniu tych zadań. Trudno dziś wyobrazić sobie spełnienie wszystkich wymogów przy zachowaniu wysokiej efektywności bez nowoczesnych systemów pomiarowych.

Umożliwiają one automatyzację procesów, pracę w chmurze oraz dostęp do danych z dowolnego miejsca. Wyniki pomiarów można od razu wysłać klientom e-mailem, co przyspiesza rozliczenia.

Wiele firm dąży dziś do wdrożenia predykcyjnego utrzymania ruchu. Jaką rolę odgrywa w tym obszarze aparatura pomiarowa i w jaki sposób firma Conrad wspiera klientów?

Predictive Maintenance to najbardziej zaawansowany obszar nowoczesnej techniki pomiarowej. Opiera się na ciągłym monitorowaniu danych, aby przewidywać awarie jeszcze przed ich wystąpieniem. Dzięki temu firmy mogą zapobiegać przestojom zamiast reagować dopiero wtedy, gdy maszyna przestanie działać. W firmie Conrad koncentrujemy się na trzech elementach wspierających ten proces. Po pierwsze, oferujemy inteligentne przyrządy pomiarowe z funkcjami rejestracji i eksportu danych. Po drugie, dostarczamy multisensory, które umożliwiają jednoczesny pomiar wielu parametrów. Po trzecie, zapewniamy rozwiązania z interfejsami programistycznymi, które automatycznie przekazują dane pomiarowe do systemów utrzymania ruchu.

Efektywność energetyczna to bardzo ważny temat. Jakie rozwiązania z zakresu techniki pomiarowej poleca Pan do pomiaru i optymalizacji zużycia energii?

Rosnące ceny energii i coraz bardziej rygorystyczne normy środowiskowe zwiększyły znaczenie nowoczesnej aparatury pomiarowej w poprawie efektywności energetycznej. Analizatory jakości energii pozwalają wykrywać



Harald Lehner, uznany ekspert w dziedzinie techniki pomiarowej, pełni funkcję kierownika kategorii marki Voltcraft w firmie Conrad Electronic (Fot. Conrad Electronic)



Technika pomiarowa do zastosowań profesjonalnych – jedna z głównych kategorii asortymentowych na Conrad Sourcing Platform (Fot. Kristof Lemp)

moc bierną i wdrażać odpowiednie działania kompensacyjne. Mierniki cęgowe pomagają identyfikować nieoczekiwane wzrosty natężenia prądu. Kamery termowizyjne umożliwiają lokalizację strat energii w budynkach oraz systemach izolacyjnych. Z kolei mierniki zużycia energii wspierają kontrolę kosztów i monitorowanie pracy mniejszych odbiorników energii.

Czy może Pan podać przykład, w którym klient Conrad osiągnął wymierne korzyści dzięki nowoczesnej aparaturze pomiarowej?

Producent maszyn dla branży farmaceutycznej regularnie miał trudności z ustaleniem przyczyn awarii swoich urządzeń. W ramach gwarancji serwisanci często musieli wyjeżdżać do klientów na wiele dni, a czasem nawet tygodni, aby zlokalizować źródło usterki i przywrócić maszynę do pracy.

W jaki sposób Conrad wsparł klienta?

Nasz przedstawiciel handlowy zabrał ze sobą kamerę akustyczną Fluke ii915 na wizytę u klienta w celu przeprowadzenia demonstracji. Klient początkowo podchodził do tego sceptycznie, jednak na miejscu udało się za pomocą tej specjalistycznej kamery szybko zlokalizować przyczynę usterki: uszkodzone łożysko w zupełnie nowym, bardzo kosztownym silniku elektrycznym. Następnie

przeprowadzono test porównawczy na sprawnym silniku. Nasz ekspert z działu wsparcia technicznego firmy Conrad zdiagnozował inne uszkodzone silniki i stwierdzono, że był to ogólny problem po stronie dostawcy silników. Korzyści dla naszego klienta: łatwa identyfikacja usterek przed wysyłką maszyn, mniej zgłoszeń gwarancyjnych, niższe koszty i zoptymalizowany proces produkcyjny w dłuższej perspektywie.



Trendy w technice pomiarowej: wywiad z Harald Lehner, ekspertem firmy Conrad (Fot. Conrad Electronic)

Informacje i sprzedaż

Conrad Electronic Sp. z o.o., <https://www.conrad.pl/>

Radioodbiorniki FM

Radio FM wydaje się tematem zamkniętym – a jednak własny odbiornik wciąż jest jednym z najwzniekszych projektów warsztatowych. Zmieniło się tylko to, czym go budujemy: zamiast cewek i żmudnego strojenia mamy dziś scalone tunery „radio na chipie”, które całą część wielkiej częstotliwości – z synteza, demodulacją i dekodowaniem RDS – zamykają w jednym układzie sterowanym magistralą I²C. Omówienie tego tematu zawiera dwie części. W części A prezentujemy projekt przykładowy – prosty odbiornik FM z RDS, niewymagający strojenia, dostępny jako zestaw AVT. W części B poddajemy go komentarzowi redakcyjnemu, omawiamy spektrum sposobów budowy odbiorników FM i zestawiamy projekty pokrewne oraz ofertę kitów AVT.



Część A. Projekt przykładowy Radioodbiornik FM z RDS (AVT5540)

Moduł odbiornika wykonano z popularnych i łatwych w montażu elementów, więc jego budowa nie sprawi trudności nawet początkującym. W użytkowaniu jest prosty dzięki interfejsowi złożonemu z wyświetlaczem LCD i dwóch impulsatorów, a dodatkowy wzmacniacz 2×1 W pozwala podłączyć głośniki bez dodatkowej elektroniki.

Najważniejsze właściwości:

- nie wymaga strojenia;
- odbiór stacji w zakresie 87,5...108 MHz;
- odbiór i wyświetlanie informacji RDS;
- pamięć 8 stacji radiowych oraz ostatnich nastaw;
- interfejs: wyświetlacz LCD 2×16 i dwa impulsatory;
- współpraca z głośnikami 4...8 Ω (wbudowany wzmacniacz 2×1 W);
- zasilanie 7...15 V DC/0,3 A; wymiary płytki 124×40 mm.

Opis układu

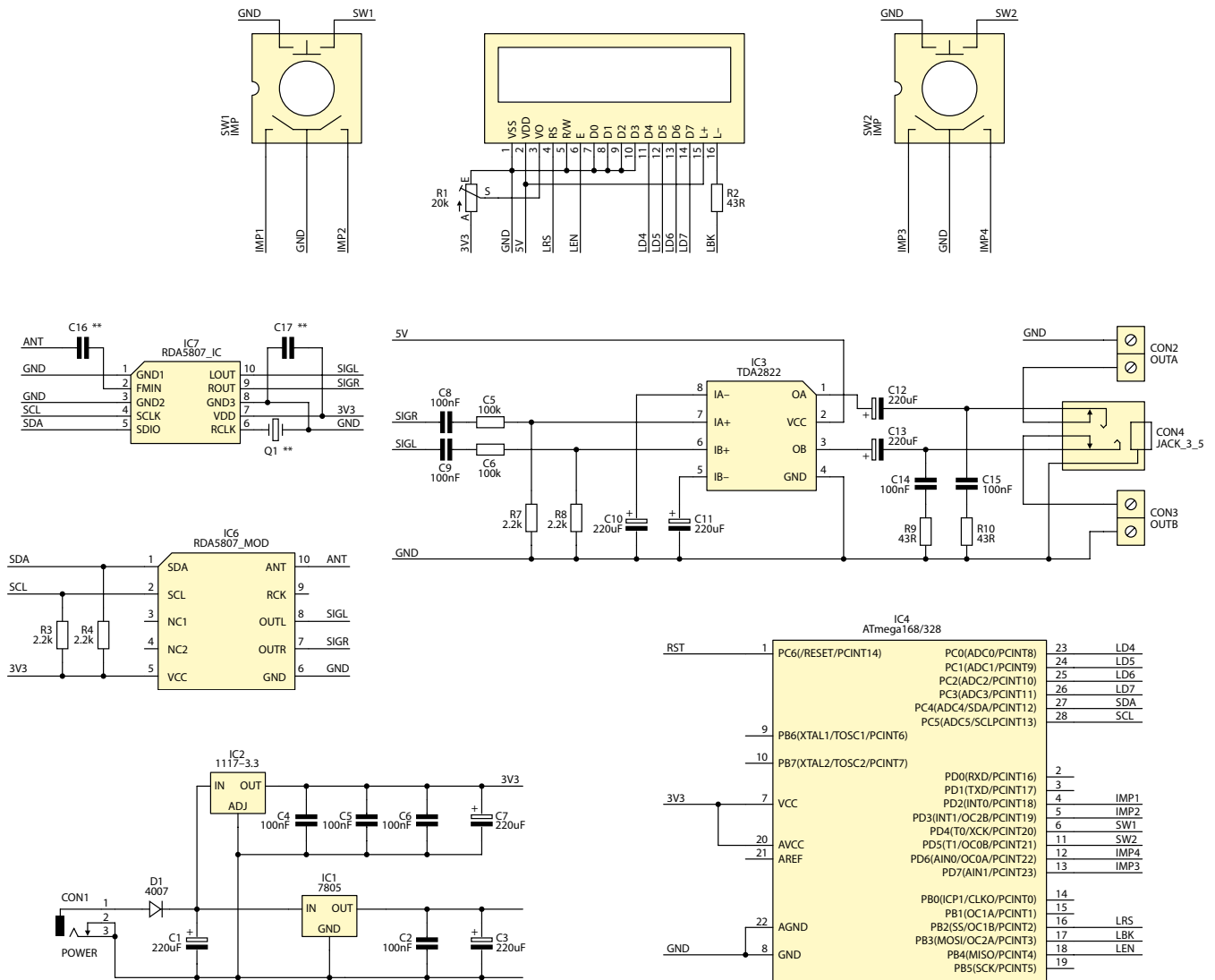
Schemat ideowy radioodbiornika pokazano na **rysunku 1**. Jego budowę można podzielić na kilka bloków: zasilania, radiowy, wzmacniacza mocy audio oraz sterowania i interfejsu użytkownika. Blok zasilania dostarcza dwóch napięć stabilizowanych: +5 V – dla wzmacniacza mocy audio i wyświetlacza, oraz +3,3 V – dla

Projekt:	radioodbiornik FM (87,5...108 MHz) z odbiorem RDS, zbudowany na scalonym module tunera RDA5807 i mikrokontrolerze ATmega168/328.
Interfejs:	wyświetlacz LCD 2×16 i dwa impulsatory; wbudowany wzmacniacz audio 2×1 W (TDA2822).
Trudność montażu:	2/5; wymiary płytki 124×40 mm.
Źródło:	instrukcja montażu zestawu AVT5540 (projekt „Radioodbiornik dla każdego”, autor: Damian Sosnowski, „Elektronika Praktyczna” 05/2016).
Zastosowanie:	prosty odbiornik FM niewymagający strojenia, z pamięcią 8 stacji i odbiorem RDS; współpracuje z głośnikami 4...8 Ω. Zasilanie 7...15 V DC/0,3 A.

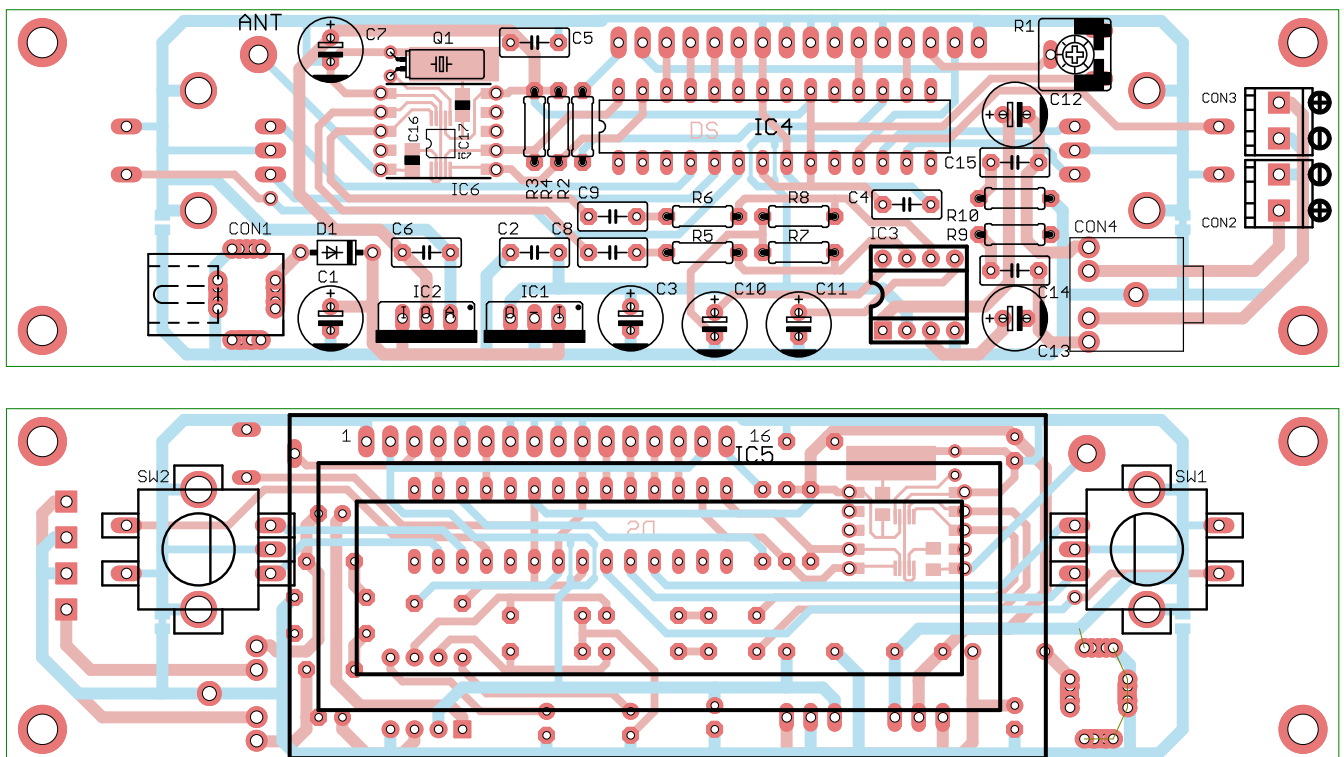
modułu radiowego i mikrokontrolera sterującego (stabilizatory IC1 7805 oraz IC2 LM1117-3.3).

Sercem części radiowej jest scalony tuner RDA5807. Układ ten ma wbudowany wzmacniacz audio małej mocy, który pozwala wysterować np. słuchawki. Aby nie obciążać tak delikatnego wyjścia i uzyskać większą moc, zastosowano dodatkowy wzmacniacz mocy audio – typową aplikację układu TDA2822 (IC3), osiągającą moc rzędu kilkuset miliwatów. Sygnał wyjściowy doprowadzono do złącza X2, co pozwala bez problemu podłączyć głośniki 4...8 Ω.

Tuner RDA5807 komunikuje się z mikrokontrolerem (IC4, ATmega168/328) poprzez interfejs I²C, a jego pracę



1. Schemat ideowy radioodbiornika FM z RDS



2. Schemat montażowy radioodbiornika FM z RDS

Tabela 1. Wykaz elementów zestawu AVT5540

Grupa	Elementy
Rezystory	PR1: potencjometr 5...20 kΩ; R2, R9, R10: 43 Ω; R3, R4, R7, R8: 2,2 kΩ; R5, R6: 100 kΩ
Kondensatory	C1, C3, C7, C10...C13: 220 μF (!); C2, C4, C5, C6, C8, C9: 100 nF; C14, C15: 100 nF
Półprzewodniki	D1: 1N4007 (!); IC1: 7805 (!); IC2: LM1117-3.3 (!); IC3: TDA2822 (!); IC4: ATmega168 lub 328 (!)
Pozostałe	SW1, SW2: impulsator z przyciskiem; IC5: wyświetlacz LCD 2×16; IC6: moduł z układem RDA5807; X1: GN DC 2.1/5.5 do druku; X2: GD381-3.5/2
Uwaga: elementy oznaczone w wykazie wykrzyknikiem (!) są biegunowe – podczas ich montażu trzeba zwrócić uwagę na orientację.	

kontroluje się przez 16 szesnastobitowych rejestrów. Mikrokontroler obsługuje wyświetlacz LCD 2×16 (IC5) oraz dwa impulsatory z przyciskami (SW1, SW2).

Montaż i uruchomienie

Rozmieszczenie elementów pokazano na rysunku 2. Montaż prowadzi się zgodnie z ogólnymi zasadami – od elementów najmniejszych do największych gabarytowo. Moduł radiowy jest już przylutowany do płytki; wyświetlacz oraz impulsatory montuje się od strony lutowania. Po zmontowaniu wystarczy ustawić kontrast wyświetlacza potencjometrem PR1 – radioodbiornik jest gotowy do pracy.

Obsługa

Na wyświetlaczu prezentowane są podstawowe informacje (rysunek 3): słupek po lewej stronie ilustruje poziom mocy odbieranego sygnału radiowego, w części centralnej widnieje aktualna częstotliwość, a słupek po prawej pokazuje poziom głośności. Osobne wskaźniki sygnalizują odbiór danych RDS, dostępność rozszerzonej informacji RDS oraz odbiór stereofoniczny.

Po kilku sekundach bezczynności, jeśli możliwy jest odbiór RDS, wskazanie częstotliwości zostaje „zasłonięte” podstawową informacją RDS (do 8 znaków – zwykle



4. Zapamiętanie ustawionej częstotliwości

nazwa stacji na zmianę z nazwą programu lub wykonawcy), a w dolnej linii pojawia się rozszerzona informacja RDS (do 64 znaków, przewijana, aby pokazać cały komunikat).

Do obsługi służą dwa impulsatory: lewy ustawia odbieraną częstotliwość, prawy reguluje głośność. Przyciśnięcie lewego impulsatora pozwala zapamiętać aktualną częstotliwość w jednej z 8 lokacji pamięci – po wybraniu numeru programu potwierdza się to ponownym przyciśnięciem (rysunek 4). Urządzenie zapamiętuje ostatnio wybrany program i głośność, a po włączeniu zasilania wznowia odbiór z tymi nastawami. Przyciśnięcie prawego impulsatora przełącza odbiór na kolejny zapisany program.

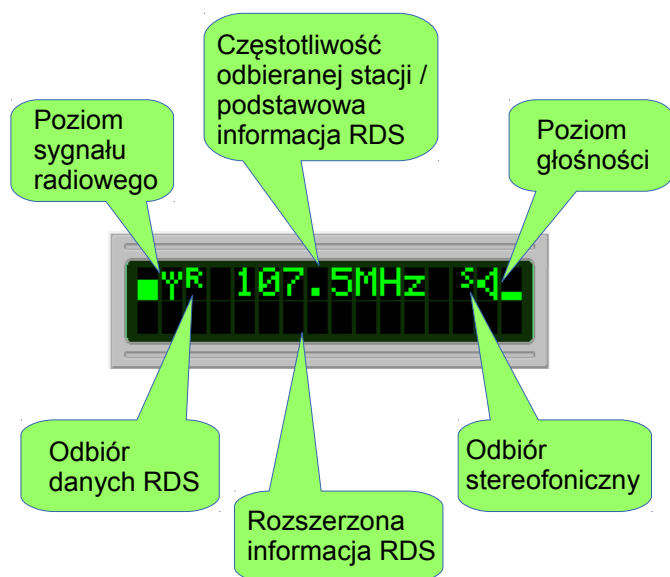
Część B. Komentarz, dyskusja rozwiązań i projekty pokrewne

Część B prowadzimy w trzech wątkach. Najpierw poddamy projekt przykładowy komentarzowi redakcyjnemu, następnie omawiamy spektrum sposobów budowy odbiorników FM, a na koniec zestawiamy projekty pokrewne, literaturę oraz ofertę kitów AVT.

Komentarz redakcyjny do projektu AVT5540

To rozsądnie pomyślany, przystępny projekt, którego mocną stroną jest właśnie prostota wynikająca z użycia scalonego tunera. Poniżej zbieramy uwagi i kierunki ewentualnej modernizacji.

„Radio na chipie” zamiast cewek – RDA5807 integruje kompletny tor FM: syntezę częstotliwości, cyfrowy demodulator i dekodery RDS, sterowane magistralą I²C. Stąd właśnie bierze się „brak strojenia”, mała płytka i niski koszt – całą trudną elektronikę



3. Informacje przedstawiane na wyświetlaczu

Tabela 2. Porównanie sposobów budowy odbiorników FM

Metoda	Zasada	Zalety	Ograniczenia
Superheterodyna	mieszanie do p.cz. 10,7 MHz	dobra czułość i selektywność	cewki, strojenie, większy układ
Scalony tuner (RDA5807/TEA5767)	synteza + demodulator + RDS, I ² C	bez strojenia, mały, tani	antena = przewód, zależność od biblioteki
Tuner DSP (Si4703/Si4735)	DSP, AM/FM/SW/SSB, RDS	wiele pasm, wysoka jakość	droższy, 3,3 V, złożona konfiguracja
SDR (np. RTL-SDR)	demodulacja w oprogramowaniu	maksymalna elastyczność	wymaga komputera/SoC, złożoność

wielkiej częstotliwości wykonano fabrycznie w jednym układzie.

Wzmacniacz TDA2822 – klasyka do wymiany – to sprawdzony, ale archaiczny i mało sprawny układ. Współcześnie naturalnym wyborem byłby wzmacniacz klasy D (np. PAM8403): z tego samego zasilania daje więcej mocy i mniej ciepła, a zajmuje mniej miejsca.

Zasilanie dwunapięciowe i straty – stabilizatory 7805 i LM1117 są liniowe; przy napięciu wejściowym dochodzącym do 15 V i poborze 0,3 A na 7805 wydzielą się zauważalna moc strat. W wersji „terenowej” lub baterijnej warto rozważyć przetwornicę impulsową.

Antena i pewność RDS – czułość i pewność dekodowania RDS zależą od anteny. Sami użytkownicy zwracają uwagę, że przydałaby się możliwość użycia przewodu (np. głośnikowego) jako anteny – drobny, lecz praktyczny kierunek rozwoju.

Wyszukiwanie stacji (seek) – RDA5807 ma sprzętową funkcję automatycznego wyszukiwania stacji. Interfejs użytkownika mógłby ją wykorzystać do skanowania pasma, zamiast wyłącznie ręcznego strojenia impulsatorem.

Jak buduje się odbiornik FM – przegląd rozwiązań

Superheterodyna (klasyczna). Sygnał miesza się do częstotliwości pośredniej 10,7 MHz, filtruje i demoduluje. Daje dobrą czułość i selektywność, ale wymaga cewek, filtrów i strojenia – droga edukacyjnie cenna, lecz kłopotliwa.

Scalony tuner „radio na chipie”. Układy takie jak RDA5807 (jak w projekcie) czy TEA5767 zamykają cały tor FM w jednym chipie z cyfrową synteza i sprzętowym RDS, sterowanym I²C. Brak strojenia, mały rozmiar, niski koszt; antena to często po prostu przewód.

Tuner DSP (Silicon Labs Si4703/Si4735). Zaawansowane układy DSP odbierają nie tylko FM, ale i AM/SW, a Si4735 także SSB; obsługują RDS. Wyższa jakość i wielopasmowość kosztem ceny, zasilania 3,3 V i bardziej złożonej konfiguracji.

Odbiornik programowy (SDR). Próbkowanie i demodulacja w oprogramowaniu (np. RTL-SDR) dają maksymalną elastyczność, ale wymagają komputera lub mocniejszego SoC i większej wiedzy. Metody zestawia tabela 2.

WM

Projekty pokrewne i literatura

Poniższe zestawienie dobrano pod kątem czytelnika zaawansowanego – obejmuje materiały o istotnej wartości warsztatowej i konstrukcyjnej; pominięto wątki forów i mediów społecznościowych.

Biblioteki i odbiorniki na scalonych tunerach

- **PU2CLR – RDA5807 Arduino Library.** biblioteka do tego samego tunera co w projekcie, z gotowymi przykładami obsługi RDS i wyszukiwania stacji. <https://github.com/pu2clr/RDA5807>
- **PU2CLR – SI470X Arduino Library (Si4702/Si4703).** tuner FM z RDS firmy Silicon Labs – alternatywa dla RDA5807. <https://github.com/pu2clr/SI470X>
- **PU2CLR – SI4735 Arduino Library.** rozbudowany odbiornik DSP AM/FM/SW/SSB z RDS; ponad 60 przykładów, popularny wśród krótkofalowców. <https://github.com/pu2clr/SI4735>

Moduł i komponent

AVT/sklep.avt.pl – moduł radio FM RDA5807M z RDS. gotowy moduł tunera (ten sam układ co w projekcie) do budowy własnego odbiornika. <https://sklep.avt.pl/modul-radio-fm-rda5807m-z-rds.html>

Radioodbiorniki i moduły FM w ofercie kitów AVT

Projekt przykładowy jest dostępny jako gotowy zestaw AVT. Poniżej zestawiamy ten kit oraz powiązane układy radiowe ze sklepu AVT (sklep.avt.pl); dostępność do potwierdzenia w sklepie.

- **AVT5540 Radio FM z RDS.** projekt przykładowy z części A – RDA5807, ATmega168/328, RDS, pamięć 8 stacji, wzmacniacz 2×1 W.
- **Moduł RDA5807M z RDS (komponent).** gotowy moduł tunera FM do własnych konstrukcji odbiorników z RDS.
- **AVT495 Miniaturowy odbiorniczek FM.** prosty, analogowy odbiornik FM – propozycja na początek.

Od superheterodyny z cewkami po „radio na chipie” sterowane magistralą I²C – odbiornik FM pozostaje projektem, w którym cała trudna elektronika wielkiej częstotliwości zmieściła się w jednym układzie, a konstruktorowi została najprzyjemniejsza część: interfejs, obudowa i dobre brzmienie.

Czujniki pojemnościowe położenia

Położenie to wielkość, którą w elektronice mierzy się na dziesiątki sposobów – potencjometrem, enkoderem optycznym, czujnikiem indukcyjnym (LVDT), układem z czujnikiem Halla, a także metodą pojemnościową. Ta ostatnia bywa najprostsza mechanicznie: dwie metalowe okładki rozdzielone powietrzem tworzą kondensator, którego pojemność zależy od odległości między nimi – wystarczy ją zmierzyć. Omówienie tego tematu zawiera dwie części. W części A prezentujemy projekt przykładowy – elegancki, ośmioelementowy czujnik położenia, który mierzy mikrometry, porównując pojemność czujnika z pojemnością wzorcową. W części B podajemy go komentarzowi redakcyjnemu, omawiamy spektrum metod pomiaru położenia i pojemności oraz zestawiamy projekty pokrewne i ofertę kitów AVT.

Część A. Projekt przykładowy Pojemnościowy czujnik położenia o rozdzielczości ~0,1%

Trudno o prostszy czujnik elektromechaniczny niż pojemnościowy: tworzą go w istocie dwie okładki (albo nawet jedna, jeśli mierzony obiekt jest przewodzący) rozdzielone dielektrykiem, na przykład powietrzem. Pojemność wynosi w przybliżeniu $C = 8,854 \text{ pF} \cdot S/d$, gdzie S to powierzchnia okładek, a d – odległość między nimi (obie wielkości w metrach). Pojemność C staje się więc czułym wskaźnikiem odległości między okładkami.

Weźmy realistyczny przykład. Okrągłe okładki o średnicy 38 mm i początkowej odległości 1 mm dają pojemność nominalną około 10 pF. Wartość ta zmienia się od 3,3 pF przy $d = 3 \text{ mm}$ do 33 pF przy $d = 0,3 \text{ mm}$ – co po prostych przeliczeniach można zamienić na odległość między okładkami. Pozostaje pytanie: jak zmierzyć tę pojemność?

Projekt: pojemnościowy czujnik położenia (przemieszczenia) mierzący odległość między okładkami z rozdzielczością bliską 0,1%. Układ interfejsu z ośmiu tanich, łatwo dostępnych elementów, oparty na poczwórnej bramce NAND z wejściami Schmitta MC74AC132.

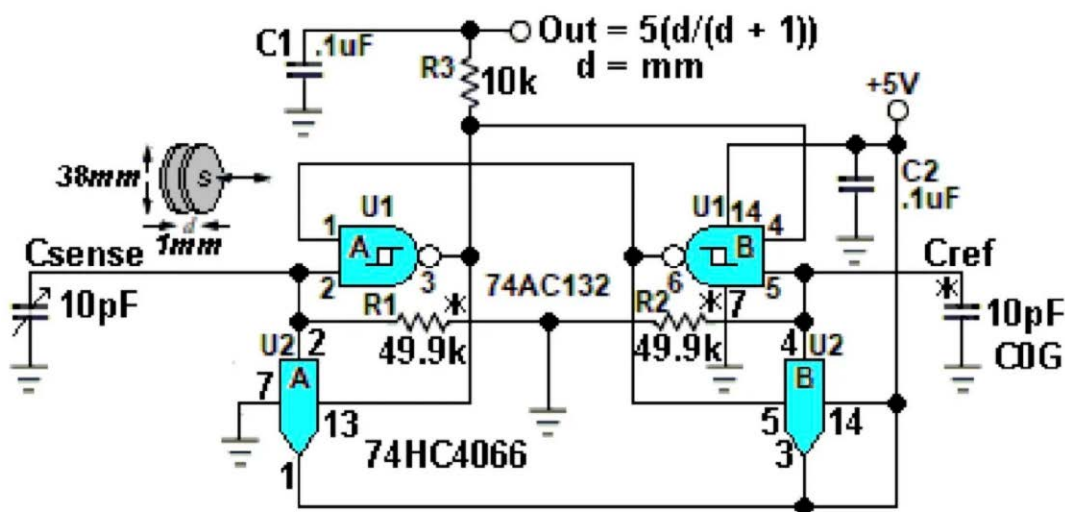
Źródło: Stephen Woodward, „~0.1% resolution capacitive position sensor”, EDN – Design Idea (27.05.2026).

Zastosowanie: precyzyjny pomiar małych przemieszczeń (rzędu mikrometrów) – stoliki pozycjonujące, czujniki nacisku i ugięcia, kontrola szczeliny; z możliwością pomiaru w dwóch osiach (XY).

Układ pomiarowy

Odpowiedzią jest prosty układ interfejsu z rysunku 1 – zaledwie osiem tanich, dostępnych „z półki” elementów – dopasowany prostotą do samego czujnika.

Para $U1a-U2a$ buduje człon czasowy RC o stałej równej $R1 \cdot C_{\text{sense}}$, a $U1b-U2b$ – analogiczny człon $R2 \cdot C_{\text{ref}}$. Po sprzężeniu krzyżowym, pokazanym



1. Dwie bramki z wejściami Schmitta (U1a, U1b) sprzężone krzyżowo tworzą multiwibrator RC o częstotliwości około 1 MHz



2. Charakterystyka przy sygnale wyjściowym podanym na 12-bitowy ADC z odniesieniem +5 V. Czarna krzywa (lewa oś) – odległość d w milimetrach; czerwona (prawa oś) – rozdzielczość najmłodszego bitu ADC w mikrometrach. W znacznej części zakresu rozdzielczość jest bliska 0,1% (d/1000)

na rysunku 1, na wyprowadzeniu 3 układu U1 powstaje generator przebiegu prostokątnego o częstotliwości rzędu 1 MHz. Przebieg ten pozostaje na poziomie +V przez czas $T_{ref} = 50 \text{ k}\Omega \cdot C_{ref} = 500 \text{ ns}$ oraz na poziomie zera przez czas $T_{sense} = 50 \text{ k}\Omega \cdot C_{sense} = 500 \text{ ns/d}$, gdzie d wyrażono w milimetrach.

Współczynnik wypełnienia sygnału wyjściowego wynosi zatem $x = T_{ref}/(T_{ref} + T_{sense}) = d/(d + 1)$. Po prostym przekształceniu otrzymujemy $x(d + 1) = d$, następnie $x = d/(1 + x)$, a ostatecznie $d = x/(1 - x)$. Aby uzyskać wynik liczbowy, sygnał wyjściowy wystarczy podać na przetwornik analogowo-cyfrowy.

Rozdzielczość i odczyt

Rysunek 2 pokazuje, jak ta zależność sprawdza się, gdy sygnał wyjściowy trafia na 12-bitowy przetwornik ADC z napięciem odniesienia +5 V.

Dokładność cyfrowego odczytu współczynnika wypełnienia wymaga, by przetwornik ADC korzystał z tego samego napięcia odniesienia +5 V, którym zasilany jest układ – wtedy oba pomiary skalują się tak samo i odczyt jest poprawny.

Uwagi praktyczne

Atutem układu jest naturalne dopasowanie i wzajemne śledzenie parametrów bramek U1 – wynikające z tego, że leżą na jednym krzemie. Elementy oznaczone gwiazdką (R1, R2 i Cref) muszą być precyzyjne, a pojemności pasożytnicze montażu i ścieżek należy starannie minimalizować.

Jeśli istnieje ryzyko, że okładki czujnika się zetkną i zewrą, warto zabezpieczyć U2 kondensatorem szeregowym 0,1 μF (z tej samej „torebki” co C1 i C2). Jest on na tyle większy od Csense, że jego tolerancja i stabilność nie mają znaczenia; dla symetrii podobny kondensator można dodać szeregowo z Cref. Gdy potrzebny jest pomiar w dwóch osiach (np. stolik XY), wolne połówki układów U1

i U2 są gotowe do wykorzystania – co pozwala zachować równie pojemną, co prostą konstrukcję.

Część B. Komentarz, dyskusja rozwiązań i projekty pokrewne

Część B prowadzimy w trzech wątkach. Najpierw poddamy projekt przykładowy komentarzowi redakcyjnemu – wskazujemy jego mocne strony oraz miejsca, w których konstruktor poprowadziłby rzecz inaczej. Następnie omawiamy spektrum metod pomiaru położenia i pomiaru pojemności. Na koniec zestawiamy projekty pokrewne, literaturę oraz ofertę kitów AVT.

Komentarz redakcyjny do projektu

To podręcznikowy przykład tego, jak z dosłownie kilku bramek wycisnąć precyzyjny pomiar. Największą zaletą jest ratiometryczność: układ porównuje Csense z Cref, więc wynik (współczynnik wypełnienia) jest w pierwszym przybliżeniu niewrażliwy na napięcie zasilania, częstotliwość generatora i dryf parametrów – pod warunkiem, że oba człony korzystają z tych samych, dobrze dopasowanych bramek na jednym krzemie. Poniżej zbieramy uwagi warte rozważenia.

Nieliniowa charakterystyka – zależność $d = x/(1 - x)$ jest z natury nieliniowa; wymaga przeliczenia w oprogramowaniu. Co eleganckie, da się ją zlinearyzować jednym mnożeniem przez stałą (wprost z zależności $C_{sense}d = \text{const}$) – tę drogę podpowiedziano w dyskusji pod oryginałem, a autor opublikował również osobną wersję z liniowym wyjściem.

Rozdzielczość a tętnienia filtru – uśrednianie sygnału wyjściowego (R3/C1) to kompromis: zbyt mała stała czasowa daje tętnienia rzędu jednego LSB, zbyt duża – powolną odpowiedź. Dla 1 MHz i 12-bitowego ADC realnie uzyskuje się około 11 bitów efektywnych przy odpowiedzi rzędu milisekund.

Pomiar cyfrowy zamiast ADC – współczynnik wypełnienia można zmierzyć także licznikiem czasu albo zliczając jedynki i zera w nadpróbkowanym strumieniu (np. przez peryferie SPI/I²S z DMA). Taki pomiar „w dziedzinie czasu” omija nieliniowości i offsety toru analogowego.

Wrażliwość na środowisko – każdy pomiar pojemnościowy reaguje na wilgotność, temperaturę i zabrudzenie dielektryka oraz na pojemności pasożytnicze. W zastosowaniach precyzyjnych konieczne są ekranowanie, krótkie połączenia i ewentualna kompensacja temperaturowa.

Jak mierzy się pojemność i położenie – przegląd rozwiązań

Oscylator RC i pomiar wypełnienia. Metoda z projektu przykładowego: stała czasowa R-C wyznacza czas

Tabela 1. Porównanie metod pomiaru położenia

Metoda	Zasada	Zalety	Ograniczenia
Pojemnościowa	zmiana pojemności $C \sim S/d$	prosta, bezstykowa, duża czułość na małe przemieszczenia	mały zakres, wrażliwość na wilgoć/zabrudzenia i C pasożytnicze
Indukcyjna/LVDT	sprzężenie magnetyczne ruchomego rdzenia	odporna, trwała, przemysłowa	wymaga wzbudzenia AC i demodulacji, większe gabaryty
Optyczna (enkodery)	siatka + fotodetektor	bardzo duża rozdzielczość i zakres	koszt, wrażliwość na zabrudzenia
Rezystancyjna	suwak po ścieżce oporowej	najprostsza i tania	stykowa, zużycie i szumy
Magnetyczna (Hall/AMR)	pole magnesu względem czujnika	bezstykowa, odporna na brud	wrażliwość na zewnętrzne pola

trwania stanów, a ich stosunek – szukaną pojemność. Najprostszaitania,wwersjiratiometrycznejodpornanadryf; ograniczeniem jest nieliniowość i wrażliwość na pojemności pasożytnicze.

Oscylator LC i przetwornik pojemność → cyfra. Nowoczesne układy, jak Texas Instruments FDC2214, mierzą pojemność przez zmianę częstotliwości rezonansowej obwodu LC. Dają rozdzielczość rzędu femtofaradów i wysoką odporność na zaburzenia EMI, z interfejsem I²C – gotowe do pomiaru zbliżenia, poziomu cieczy i przemieszczenia.

Przetwornik $\Sigma\Delta$ pojemność → cyfra. Układy Analog Devices AD7745/AD7746 mierzą pojemność bezpośrednio,

z rozdzielczością 24-bitową i liniowością $\pm 0,01\%$, przy zakresie wejściowym ± 4 pF. Stosuje się je m.in. w sejsmometrach i czujnikach ugięcia, gdzie liczy się skrajnie mała zmiana pojemności.

Inne metody pomiaru położenia. Poza pojemnościową stosuje się metodę indukcyjną (LVDT), optyczną (enkodery), rezystancyjną (potencjometry) oraz magnetyczną (czujniki Halla, AMR). Każda ma swój kompromis między zakresem, rozdzielczością a odpornością na warunki pracy – zestawia je tabela 1.

WM

Projekty pokrewne i literatura

Poniższe zestawienie dobrano pod kątem czytelnika zaawansowanego – obejmuje materiały o istotnej wartości warsztatowej i konstrukcyjnej; pominięto wątki forów i mediów społecznościowych.

Ten sam autor (Stephen Woodward, EDN – Design Ideas)

- **EDN** – Capacitive position sensor with linearized output. rozwinięcie projektu przykładowego o liniowe wyjście (linearyzacja zależności d od współczynnika wypełnienia). <https://www.edn.com/capacitive-position-sensor-with-linearized-output/>
- **EDN** – CMOS flip-flop used “off label” implements precision capacitance sensor. pomiar pojemności z użyciem przerzutnika CMOS w nietypowej roli. <https://www.edn.com/cmos-flip-flop-used-off-label-implements-precision-capacitance-sensor/>
- **EDN** – From gap to signal: Non-contact capacitive displacement sensors. wprowadzenie do bezstykowego pomiaru przemieszczenia metodą pojemnościową. <https://www.edn.com/from-gap-to-signal-non-contact-capacitive-displacement-sensors/>
- **EDN** – Simple circuit interfaces differential capacitance sensor. interfejs do różnicowego czujnika pojemnościowego. <https://www.edn.com/simple-circuit-interfaces-differential-capacitance-sensor/>
- **EDN** – Capacitor matchmaker. praktyczna metoda dobierania i porównywania kondensatorów – przydatna przy precyzyjnym Cref. <https://www.edn.com/capacitor-matchmaker/>

Przetworniki pojemność → cyfra (noty i moduły)

- **Analog Devices** – AD7745/AD7746, 24-bitowy przetwornik pojemność → cyfra. nota katalogowa: bezpośredni pomiar C, liniowość $\pm 0,01\%$, interfejs I²C. <https://www.analog.com/en/products/ad7746.html>
- **Texas Instruments** – FDC2214, 28-bitowy przetwornik pojemność → cyfra z rezonansem LC. do pomiaru zbliżenia, poziomu cieczy i przemieszczenia. <https://www.ti.com/product/FDC2214>; moduł open hardware (KiCad): <https://protocentral.com/product/protocentral-fdc2214-capacitance-sensor-breakout/>
- **Hackaday.io** – 22-bit Capacitance to Digital Converter. domowy sejsmometr z różnicowym czujnikiem pojemnościowym i układem AD7745. <https://hackaday.io/project/165806-22-bit-capacitance-to-digital-converter>

Pomiar pojemności w ofercie kitów AVT

Projekt przykładowy nie jest zestawem AVT, ale jego sercem jest precyzyjny pomiar pojemności – a takie przyrządy są w ofercie AVT obecne. Poniżej zestawiamy najbliższe tematycznie kity ze sklepu AVT (sklep.avt.pl).

- **AVT2725 Mikroprocesorowy miernik pojemności.** szeroki zakres pomiarowy, możliwość kalibracji nawet kilkumetrowych przewodów pomiarowych – istotne przy pomiarze małych pojemności oddalonych od przyrządu.
- **AVT2425 Miernik pojemności (EdW).** prosty przyrząd warsztatowy do pomiaru pojemności.
- Od prostego multiwibratora RC po 24-bitowe przetworniki pojemność → cyfra – pomiar pojemnościowy raz po raz pokazuje, jak wiele precyzji można wydobyć z dwóch kawałków metalu i odrobiny inżynierskiej pomysłowości

REKLAMA

facebook.com/ElektronikaPraktyczna

Sterowniki wentylatorów

Dwuwentylatorowy sterownik PWM z izolowaną przetwornicą flyback 6 W

Wentylator to najprostszy, a zarazem najskuteczniejszy sposób odprowadzenia ciepła z wnętrza urządzenia. Sztuka nie polega jednak na tym, by go po prostu włączyć, lecz by kręcił się dokładnie tak szybko, jak trzeba – ani za wolno (przegrzanie), ani za szybko (hałas i zużycie). Sposobów na to jest wiele: od progowego włącz/wyłącz, przez analogową regulację napięcia, po sterowanie PWM zależne od temperatury i wentylatory czteroprzewodowe ze sprzężeniem tachometrycznym. Omówienie tego tematu zawiera dwie części. W części A prezentujemy projekt przykładowy – kompletny, dwuwentylatorowy sterownik PWM z własnym zasilaczem sieciowym. W części B poddajemy ten projekt komentarzowi redakcyjnemu, omawiamy spektrum rozwiązań spotykanych w sterownikach wentylatorów i zestawiamy projekty pokrewne oraz ofertę kitów AVT.

Część A. Projekt przykładowy Dwuwentylatorowy sterownik PWM z własnym zasilaczem sieciowym.

Przedstawiony projekt łączy w jednej, zwartej płytce dwie funkcje, które zwykle realizuje się osobno: zasilanie sieciowe i inteligentne sterowanie chłodzeniem. Pierwsza część to izolowana przetwornica zaporowa (flyback) o mocy 6 W, przekształcająca napięcie sieciowe z szerokiego zakresu 85...260 V AC na stabilne 12 V DC. Druga to układ sterujący na mikrokontrolerze ATtiny13, który odczytuje temperaturę z termistora NTC 10 kΩ i na bieżąco reguluje prędkość wentylatora sygnałem PWM – zapewniając skuteczne chłodzenie przy ograniczeniu hałasu i poboru mocy. Szczególną uwagę poświęcono obwodowi drukowanemu pod kątem zgodności z wymaganiami EMC/EMI.

Najważniejsze właściwości

- uniwersalne wejście zasilania 85...260 V AC (praca w sieci 230 V);
- izolowane galwanicznie wyjście 12 V DC o mocy 6 W;
- sterowanie jednym lub dwoma wentylatorami 12 V;
- płynna regulacja obrotów sygnałem PWM o częstotliwości 25 kHz (ponad pasmem słyszalnym);
- pomiar temperatury termistorem NTC 10 kΩ (złącze XH-2P);
- sterownik przetwornicy DK106 z tranzystorem MOSFET w strukturze – bez rezystora rozruchowego i uzwojenia pomocniczego;
- obwód drukowany zaprojektowany pod kątem EMC/EMI;
- mikrokontroler ATtiny13 programowany przez złącze ISP.

Projekt:	dwuwentylatorowy sterownik chłodzenia PWM z izolowaną przetwornicą flyback 6 W; uniwersalne wejście 85...260 V AC, wyjście 12 V DC. Sterowanie oparte na mikrokontrolerze ATtiny13, układ zasilacza na sterowniku DK106.
Źródło:	Hesam Moshiri, „Dual-FAN PWM Cooling Circuit with a 6 W Isolated Flyback Power Supply”, Power Electronics News (2025). Pliki Gerber i kod do pobrania z materiału źródłowego.
Zastosowanie:	samodzielny moduł zarządzania termicznego w obudowach urządzeń, zasilaczach i innych miejscach wrażliwych na przegrzanie; obsługa jednego lub dwóch wentylatorów 12 V.

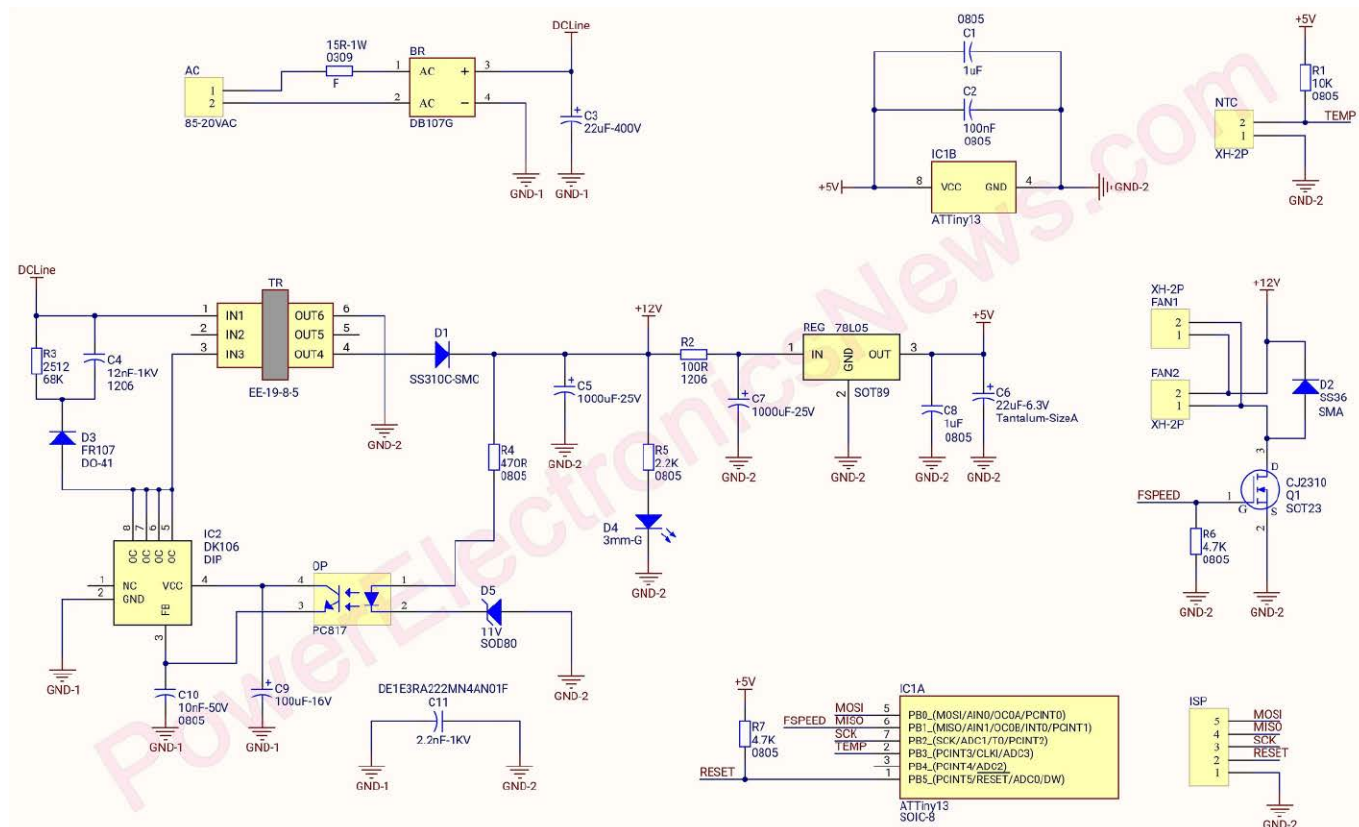
Zasada działania

Schemat ideowy urządzenia przedstawia **rysunek 1**. Składa się on z dwóch części: przetwornicy flyback AC/DC oraz sterownika wentylatora. Najważniejsze parametry układu zestawiono w **tabeli 1**.

Po stronie sterowania mikrokontroler odczytuje napięcie z dzielnika z termistorem NTC za pomocą wbudowanego przetwornika ADC. Wzrost temperatury powoduje spadek rezystancji termistora, a tym samym spadek napięcia na wejściu ADC; program przelicza tę wartość na wypełnienie sygnału PWM i steruje wentylatorem. Progi za działania oraz prędkości można dostosować w programie.

Zasilacz – przetwornica flyback

Zaciski „AC” stanowią wejście napięcia przemienicznego 85...260 V. Element F to rezystor bezpiecznikowy 15 Ω/1 W w obudowie 0309 MELF; zamiast zwykłego bezpiecznika ogranicza on prąd udarowy (włączeniowy), pełni funkcję zabezpieczenia oraz tworzy filtr RC wraz z kondensatorem C3. BR to mostek prostowniczy



1. Schemat ideowy sterownika: przetwornica flyback (po lewej) i sterownik wentylatora (po prawej), rozdzielone barierą izolacji galwanicznej

(DB107G), a C3 – kondensator szyny wysokiego napięcia (sieć DCLine) ograniczający tętnienia. W tym miejscu należy zastosować kondensator wysokiej jakości; element gorszej klasy lub o zbyt małej dopuszczalnej wartości prądu tętnień zbyt szybko wyschnie i może uszkodzić sterownik.

Elementy R3, D3 i C4 tworzą obwód tłumiący RCD (snubber), który wygasa wysokonapięciowe szpilki i oscylacje (ringing) na uzwojeniu pierwotnym transformatora (drenie tranzystora MOSFET), chroniąc tranzystor i ograniczając poziom EMI. W opisywanym układzie tranzystor MOSFET jest zintegrowany w strukturze sterownika IC2 (DK106). Sterownik pracuje z częstotliwością 65 kHz i wymaga zaledwie kilku elementów biernych – nie potrzebuje rezystora rozruchowego ani uzwojenia pomocniczego.

Kondensator C10 ogranicza zaburzenia wielkiej częstotliwości na wyprowadzeniu sprzężenia zwrotnego, a C9 jest kondensatorem blokującym wyprowadzenia VCC. Transoptor PC817 (OP) zapewnia izolowaną galwanicznie drogę przekazania sterownikowi informacji o zmianach napięcia wyjściowego. D1 to wyjściowa dioda Schottky'ego prostująca napięcie po stronie wtórnej, a C5 ogranicza zaburzenia na wyjściu. Zaświecenie diody D4 sygnalizuje poprawną wartość napięcia wyjściowego.

Sterownik wentylatora

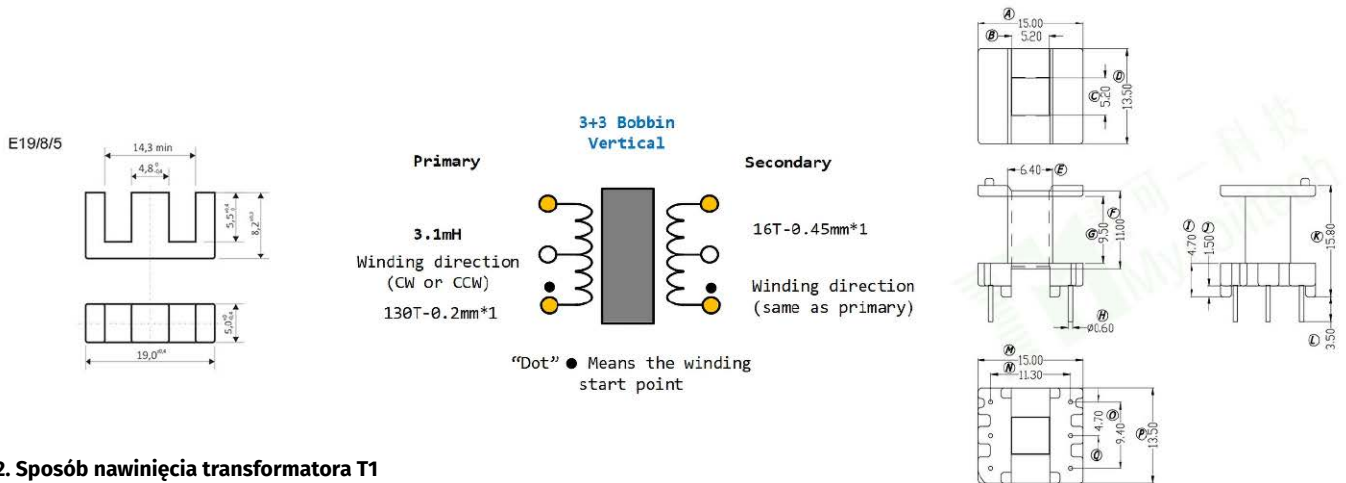
Filtr dolnoprzepustowy RC (R2, C7) maksymalnie ogranicza zaburzenia na wejściu pomiarowym

Tabela 1. Najważniejsze parametry sterownika

Parametr	Wartość
Napięcie wejściowe	85...260 V AC
Wyjście	12 V DC/6 W
Sterownik przetwornicy	DK106 (MOSFET zintegrowany), 65 kHz
Mikrokontroler	ATtiny13 (programowanie ISP)
Częstotliwość PWM wentylatora	25 kHz
Czujnik temperatury	NTC 10 kΩ (złącze XH-2P)
Transformator	EE19-8-5, karkas 3+3, Lp ≈ 3,1 mH

– przetwornik ADC mikrokontrolera jest wrażliwy na szumy. Stabilizator 78L05 w obudowie SOT-89 (REG) dostarcza napięcie +5 V dla mikrokontrolera (IC1); kondensatory C8 i C6 stabilizują jego pracę, przy czym C6 jest typu tantalowego, lepiej sprawdzającego się na wyjściu stabilizatora LDO niż ceramiczny.

IC1 to mikrokontroler ATtiny13, programowany przez 5-stykowe złącze ISP 2,54 mm w standardzie AVR. C1 i C2 to kondensatory odsprężające szyny zasilania. Tranzystor MOSFET z kanałem typu N – Q1 (SOT-23) – steruje wentylatorami po stronie dolnej (low-side), a dioda D2 chroni go przed szpilkami napięcia o przeciwnej polaryzacji. D2 musi być diodą Schottky'ego zbliżoną do SS36 – przy częstotliwości PWM 25 kHz zwykła dioda (np. M7) nie przełączałaby się dostatecznie szybko. Termistor NTC 10 kΩ dołącza się do złącza XH-2P.



2. Sposób nawinięcia transformatora T1

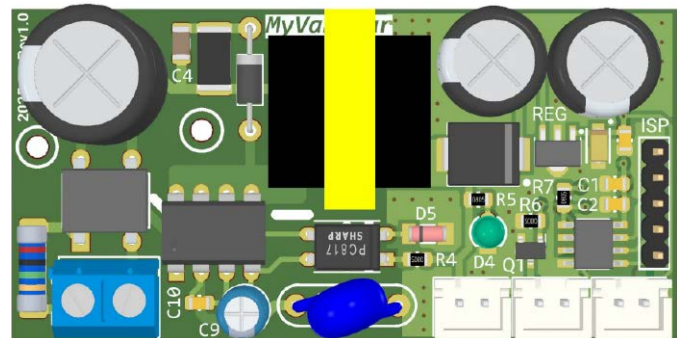
Transformator

Dane uzwojeń przedstawia rysunek 2. Rdzeń ferrytowy to EE19-8-5, a karkas – 3+3 (pionowy).

Najpierw nawija się uzwojenie pierwotne (zgodnie lub przeciwnie do ruchu wskazówek zegara). Po osadzeniu rdzenia w karkasie mierzy się indukcyjność uzwojenia pierwotnego miernikiem LCR – najlepiej przy częstotliwości 65 kHz (częstotliwość przełączania), ewentualnie 40 kHz lub 100 kHz. Szlifując środkową kolumnę rdzenia, dochodzi się do wartości bliskiej 3,1 mH (niewielka tolerancja jest dopuszczalna). Następnie owija się uzwojenie pierwotne 2...3 zwojami taśmy transformatorowej i nawija uzwojenie wtórne. Indukcyjność rozproszenia uzwojenia pierwotnego mierzy się przy zwartych pozostałych uzwojeniach.

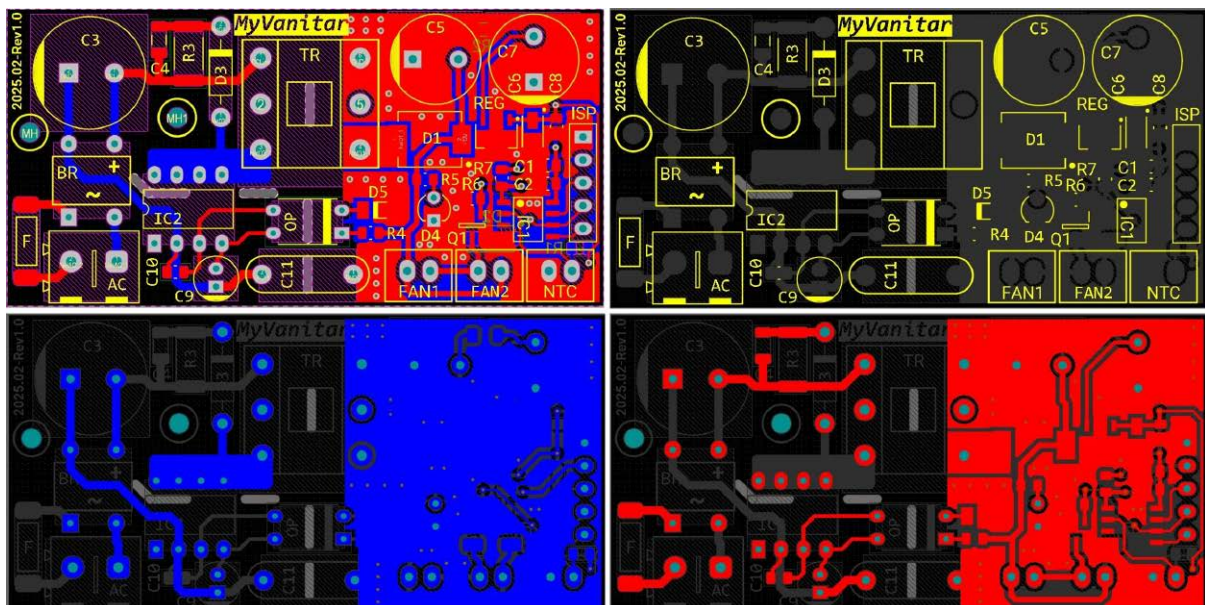
Obwód drukowany

Zastosowano dwuwarstwową płytkę z elementami SMD i przewlekkanymi (THT). Projekt przedstawia rysunek 3, a rysunek montażowy – rysunek 4.

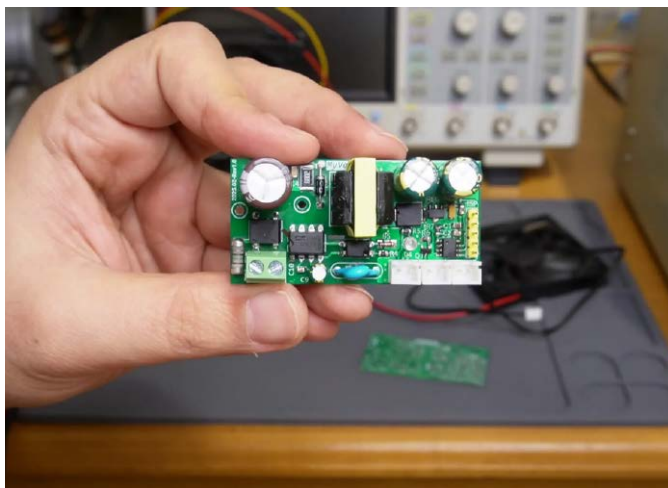


4. Rysunek montażowy płytki

Prowadzenie ścieżek podporządkowano zgodności EMC: ścieżka od wyprowadzenia drenu do transformatora jest możliwie najkrótsza; kondensator blokujący C9 leży blisko wyprowadzenia VCC; obwód tłumiący umieszczono jak najbliżej transformatora; ścieżki do kondensatora klasy Y są najkrótsze z możliwych; pole pętli po stronie wtórnej (dioda Schottky'ego i sąsiadujący z masą kondensator) jest małe.



3. Projekt obwodu drukowanego (strona górna i dolna)



5. Zmontowany sterownik

Program

Sygnał PWM jest konfigurowany poprzez rejestr mikrokontrolera AVR; jego częstotliwość wynosi 25 kHz. Pliki programu oraz pliki Gerber można pobrać z materiału źródłowego (patrz „Projekty pokrewne i literatura” w części B).

Montaż i uruchomienie

Testowanie urządzenia jest proste: wystarczy dołączyć termistor NTC 10 kΩ do złącza i ogrzać czujnik (np. opalarką lub suszarką do włosów) – wentylator powinien przyspieszać wraz ze wzrostem temperatury. Progi zadziałania i prędkość wentylatora można zmodyfikować w programie. Zmontowaną płytkę przedstawia fotografia 5.

Uwaga – bezpieczeństwo. To układ z bezpośrednim zasilaniem sieciowym i nieizolowaną stroną pierwotną. Mimo galwanicznej izolacji wyjścia uruchamianie i pomiary po stronie pierwotnej wymagają transformatora separacyjnego oraz zachowania ostrożności.

Część B. Komentarz, dyskusja rozwiązań i projekty pokrewne

Część B prowadzimy w trzech wątkach. Najpierw poddamy projekt przykładowy komentarzowi redakcyjnemu – wskazujemy jego mocne strony oraz miejsca, w których współczesny konstruktor poprowadziłby rzecz inaczej. Następnie omawiamy spektrum rozwiązań spotykanych w sterownikach wentylatorów – od progowego włącz/wyłącz po wentylatory czteroprzewodowe. Na koniec zestawiamy projekty pokrewne, literaturę oraz ofertę kitów AVT.

Komentarz redakcyjny do projektu

Projekt jest przemyślany i kompletny. Jego mocną stroną jest połączenie nowoczesnej, mocno zintegrowanej

przetwornicy flyback (DK106 z tranzystorem MOSFET w strukturze) z tanim, klasycznym sterowaniem na ATtiny13 – całość mieści się na jednej, zwartej płytce. Poniżej zbieramy uwagi warte rozważenia przy ewentualnej modernizacji.

Snubber RCD jako element obowiązkowy – w przetwornicy flyback to on decyduje, czy MOSFET przeżyje szpilki napięcia z indukcyjności rozproszenia. Autor słusznie umieszcza go tuż przy transformatorze; to nie jest miejsce na oszczędności.

Częstotliwość 25 kHz to wybór akustyczny – sterowanie PWM ponad pasmem słyszalnym eliminuje „buczenie” wentylatora przy małym wypełnieniu. Stąd wynika wymóg szybkiej diody Schottky’ego zamiast zwykłej diody prostowniczej.

Filtr RC przed wejściem ADC – pomiar temperatury miesza się z impulsami przetwornicy; bez filtra R2/C7 odczyt NTC byłby zaszumiony, a progi „pływałyby”. To detal, który odróżnia układ działający stabilnie od kapryśnego.

Layout pod EMC – krótkie pętle strony wtórnej, kondensator klasy Y przy najkrótszych ścieżkach i kondensator blokujący VCC przy wyprowadzeniu to elementarz, który w praktyce odróżnia płytkę „przechodzącą badania” od „grzejącej eter”.

Pole do rozbudowy – warto rozważyć histerezę progów w programie (przeciw „polowaniu” wentylatora przy progu), zagwarantowanie minimalnej prędkości startowej (część wentylatorów nie rusza poniżej pewnego wypełnienia PWM), dwa niezależne kanały sterowania oraz wariant dla wentylatorów czteroprzewodowych (wejście PWM + wyjście tachometryczne).

Uwaga o oznaczeniu MOSFET-a – w materiale, na którym oparto część A, tranzystor Q1 oznaczono jako Si2310 (UMW), natomiast w pierwotnej publikacji w Power Electronics News figuruje CJ2310. Oba to tranzystory N-MOSFET w obudowie SOT-23 o zbliżonych parametrach; przy zakupie warto zweryfikować dostępność i zamienniki.

Jak steruje się obrotami wentylatora – przegląd rozwiązań

Sterowanie progowe (włącz/wyłącz). Najprostsze: komparator lub mikrokontroler załącza wentylator po przekroczeniu progu temperatury (przełącznikiem lub tranzystorem). Tanie i niezawodne, lecz hałaśliwe przy załączeniu i pozbawione płynności – wentylator pracuje skokowo.

Regulacja liniowa (zmiana napięcia). Tranzystor szeregowy obniża napięcie zasilania wentylatora proporcjonalnie do temperatury. Praca jest cicha i pozbawiona zaburzeń EMI, ale tranzystor rozprasza moc i wymaga radiatora; sprawność spada przy małych obrotach.

PWM po stronie zasilania (2-/3-przewodowe). Rozwiązanie z projektu przykładowego: klucz N-MOSFET

Tabela 2. Porównanie metod sterowania obrotami wentylatorów

Metoda	Zasada	Zalety	Ograniczenia
Progową (ON/OFF)	komparator/MCU + przekaźnik lub tranzystor	prosta, tania, niezawodna	skokowa praca, hałas przy załączeniu
Liniową (napięciową)	tranzystor szeregowy	cicha, brak EMI	straty mocy, radiator, niska sprawność
PWM 2-/3-przew.	klucz N-MOSFET + dioda Schottky'ego	wysoka sprawność, płynność	wymaga $f > 20$ kHz, dobór diody
PWM 4-przewodowe	sygnał PWM do wentylatora + tacho	standard serwerowy, kontrola obrotów	wymaga wentylatora 4-przew.

włącza i wyłącza zasilanie wentylatora z dużą częstotliwością, a dioda Schottky'ego przejmuje prąd cewki. Wysoka sprawność i płynna regulacja; warunkiem braku hałasu jest częstotliwość ponad pasmem słyszalnym (tu 25 kHz).

PWM czteroprzewodowe (wejście PWM wentylatora). Wentylatory serwerowe i komputerowe mają osobne wejście sterujące PWM oraz wyjście tachometryczne (pomiar obrotów). Sterownik nie przełącza zasilania, lecz podaje

sygnał sterujący – co pozwala na pętlę regulacji z kontrolą rzeczywistych obrotów.

Drugą osią wyboru jest czujnik temperatury: termistor NTC (jak w projekcie) jest najtańszy i sprawdzony; układy analogowe typu LM35 dają napięcie proporcjonalne do temperatury; czujniki cyfrowe (DS18B20, DHT) upraszczają program kosztem dodatkowej magistrali. Metody sterowania zestawia tabela 2. **WM**

Projekty pokrewne i literatura

Poniższe zestawienie dobrano pod kątem Czytelnika zaawansowanego – obejmuje materiały o istotnej wartości warsztatowej i konstrukcyjnej; pominięto wątki forów i mediów społecznościowych.

Ten sam autor (Hesam Moshiri)

- **Power Electronics News** – Dual-FAN PWM Cooling Circuit with a 6W Isolated Flyback Power Supply. oryginał projektu przykładowego – komplet rysunków (schemat, PCB, montaż, fotografie) oraz pliki Gerber i kod. <https://www.powerselectronicsnews.com/dual-fan-pwm-cooling-circuit-with-a-6w-isolated-flyback-power-supply/>
- **EDN/Hackster** – PWM Cooling-FAN Control and Over-Temperature Protection (LM35 + ATtiny13). poprzednik bez zasilacza flyback, ze sterowaniem 3-/4-przewodowym, zabezpieczeniem nadtemperaturowym (przekaźnik) i sygnalizacją. <https://www.edn.com/pwm-cooling-fan-controller-and-over-temperature-protection-using-lm35-and-attiny13/>; <https://www.hackster.io/hesam-moshiri/3-wires-4-wires-fan-control-and-over-temperature-protection-738932>

Pomiar temperatury i sterowanie wentylatorem

- **Ametherm** – Designing a Furnace Fan and Limit Controller Using an NTC Thermistor and an Arduino. ujęcie grzewczo-przemysłowe (HVAC): NTC + Arduino Nano + przekaźniki. <https://www.ametherm.com/blog/thermistors/designing-a-furnace-fan-and-limit-controller-using-an-ntc-thermistor-and-an-arduino-microcontroller/>
- **Homemade Circuits** – Temperature Controlled DC Fan Circuits. przegląd od progowego ON/OFF po liniowe PWM z LM35/DHT11 i MOSFET – materiał wprowadzający. <https://www.homemade-circuits.com/automatic-temperature-regulator-circuit/>
- **Medium (M. Malcolm)** – PWM Fan Controller for Dummies – using an Arduino. dzielnik z NTC 10 kΩ, linearyzacja, projekt PCB w KiCad – dobra podbudowa o pomiarze NTC. https://medium.com/@malcolmmalcolm_19/pwm-fan-controller-for-dummies-using-an-arduino-cdce791a183e

Sterowniki mikrokontrolerowe i warianty

- **DPS3005C (blog)** – ATtiny13 multi-speed fan controller. NTC + ATtiny13 + MOSFET; cenne uwagi o minimalnym wypełnieniu startowym i doborze RDS(on). <https://dps3005cfancontroller.wordpress.com/>
- **Hackaday** – An ATtiny GPU Fan Controller That Sticks. ATtiny85, sterowanie PWM wentylatora GPU; praktyczna dyskusja o sygnale tachometrycznym. <https://hackaday.com/2025/08/02/an-attiny-gpu-fan-controller-that-sticks/>
- **Hackaday.io** – Arduino Variable Speed Temp. Fan Controller. termistor 10 kΩ + Arduino Nano, liniowe narastanie PWM – przykład chłodzenia wzmacniacza. <https://hackaday.io/project/158868-arduino-variable-speed-temp-fan-controller>

Sterowniki wentylatorów w ofercie kitów AVT

Projekt przykładowy nie jest zestawem AVT, ale tematyka sterowania wentylatorami jest w ofercie AVT szeroko reprezentowana. Poniżej zestawiamy aktualne kity ze sklepu AVT (sklep.avt.pl).

Regulatory obrotów wentylatorów

- **AVT478 Regulator obrotów dwóch wentylatorów 12 V.** niezależne sterowanie dwoma wentylatorami zależnie od temperatury chłodzonego urządzenia; do 10 W.
- **AVT1596 Regulator obrotów wentylatora (PC).** dopasowuje obroty do zapotrzebowania na chłodzenie; 12 V DC.
- **AVT3243 Analogowy szeregowy regulator obrotów.** regulacja zależna od temperatury radiatora – rozwiązanie „bez procesora”.
- **AVT1564 Sterownik wentylatora 12 V.** wariant dla pojedynczego wentylatora.
- **AVT1613 Regulator obrotów wentylatora 230 V.** dla wentylatorów z małymi silnikami indukcyjnymi zasilanymi z sieci 230 V AC.

Termostaty i regulatory temperatury

- **AVT5790 Stabilizator temperatury do szafki RTV.** gotowy moduł utrzymujący zadaną temperaturę.
- **AVT1699 Regulator temperatury.** uniwersalny „mózg” termiczny zamiast NTC + MCU.

Moduły zasilające

- **AVT1895/12 Uniwersalny moduł zasilający 12 V.** odpowiednik bloku zasilania z projektu, gdyby rozdzielić zasilacz od sterownika.

Od progowego włącz/wyłącz po wentylatory czteroprzewodowe ze sprzężeniem tachometrycznym – sterowanie wentylatorem pozostaje wdzięcznym tematem warsztatowym, w którym prostota układu idzie w parze z realną poprawą komfortu i trwałości chłodzonego urządzenia.

Zgrzewarki oporowe

Zgrzewarka oporowa to jeden z tych przyrządów, które elektronik najczęściej buduje sam – zwłaszcza odkąd modą stało się składanie pakietów z ogniw 18650. Niemal wszystkie amatorskie konstrukcje bazują na przerobionym transformatorze od kuchenki mikrofalowej: tani, łatwo dostępny ze złomu, a po przezwojeniu z wysokiego napięcia na niskie i wieloprądowe daje moc zbliżoną do 1 kW. Diabeł tkwi jednak w sterowaniu – w momencie załączenia i w doborze elementu wykonawczego. Omówienie tego tematu zawiera dwie części. W części A prezentujemy projekt przykładowy – sprawdzony, uniwersalny sterownik zgrzewarki, dostępny jako zestaw AVT. W części B poddajemy go komentarzowi redakcyjnemu, omawiamy spektrum rozwiązań stosowanych w zgrzewarkach i zestawiamy projekty pokrewne oraz ofertę kitów AVT.

Część A. Projekt przykładowy Sterownik zgrzewarki oporowej (AVT5553)

W sieci znajdziemy mnóstwo amatorskich zgrzewarek oporowych do łączenia drobnych elementów metalowych i mocowania blaszek (tabów) do akumulatorów. Wszystkie bazują na transformatorze z mikrofalówki, który łatwo i tanio przerabia się z wersji wysokonapięciowej na niskonapięciową i wieloprądową. Szczegółów samej przeróbki najlepiej szukać w sieci (zdjęcia, filmy); tutaj proponujemy sprawdzony sterownik do takiej zgrzewarki, minimalizujący prąd rozruchowy, pozwalający regulować czas i opóźnienie zadziałania oraz gwarantujący symetryczne zasilanie transformatora. Wiele dotychczasowych sterowników to mniej lub bardziej wierne kopie starej koncepcji – i niestety powielają jej poważne wady. Proponowany układ ma stopień wyjściowy dopasowany do obciążenia indukcyjnego i załącza je w maksimum chwilowego napięcia sieci, a więc całkowicie odwrotnie niż ponad połowa podobnych projektów. W praktyce może współpracować z dowolnym transformatorem średniej mocy.

Najważniejsze właściwości

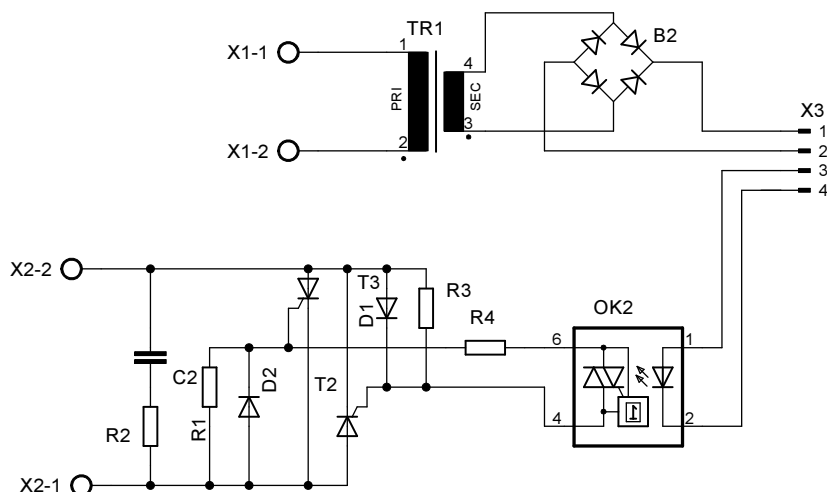
- idealny do zgrzewania pakietów akumulatorów;
- interfejs użytkownika: wyświetlacz LCD 2×16 i przyciski;
- minimalizacja prądu rozruchowego, regulacja czasu i opóźnienia zadziałania oraz symetryczne zasilanie transformatora;
- stopień wyjściowy dopasowany do obciążenia indukcyjnego – załączanie w maksimum chwilowego napięcia sieci;
- wymiary płytek: sterownika 96×50 mm, wykonawczej 96×45 mm;

Projekt:	uniwersalny sterownik zgrzewarki oporowej zbudowanej na przerobionym transformatorze od kuchenki mikrofalowej (moc ~1 kW). Mikrokontroler ATmega8, interfejs LCD 2×16 i przyciski; konstrukcja na dwóch płytkach – sterownika (96×50 mm) i wykonawczej (96×45 mm).
Źródło:	instrukcja montażu zestawu AVT5553. Zestaw powstał na podstawie projektu o tym samym tytule opublikowanego w „Elektronice Praktycznej” 10/2016. Autor: Robert Magdziak.
Rekomendacje:	idealny do zgrzewania pakietów akumulatorów i mocowania końcówek do ogniw – szczególnie cenny dla modelarzy i osób zainteresowanych energią odnawialną.

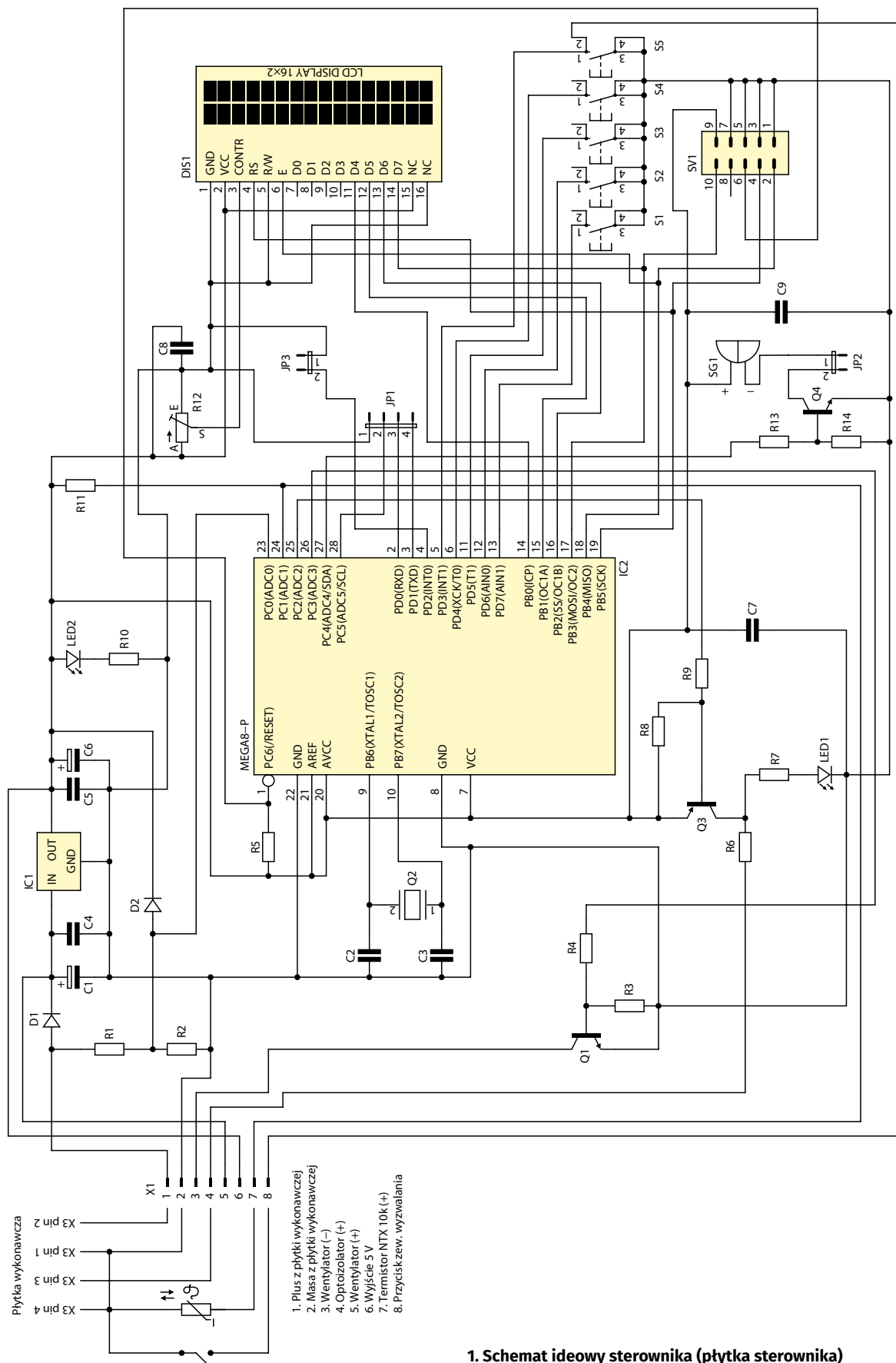
- kompletna zgrzewarka wymaga transformatora od kuchenki mikrofalowej.

Opis układu

Układ sterownika jest zasilany wyprostowanym, ale nieodfiltrowanym napięciem z transformatora sieciowego – dzięki temu może wykrywać przejścia sinusoidy sieci przez zero. Służy do tego dzielnik R1/R2 z diodą



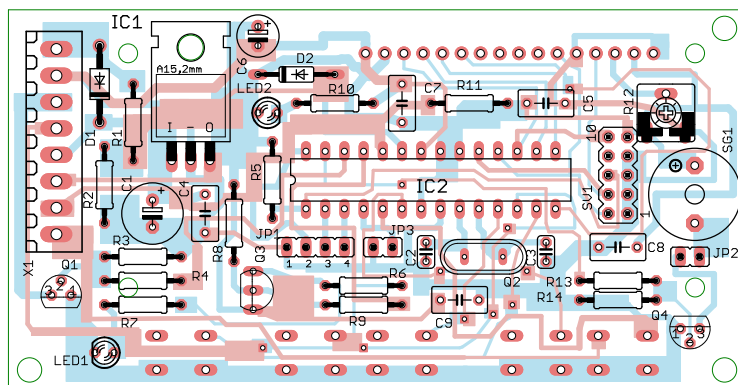
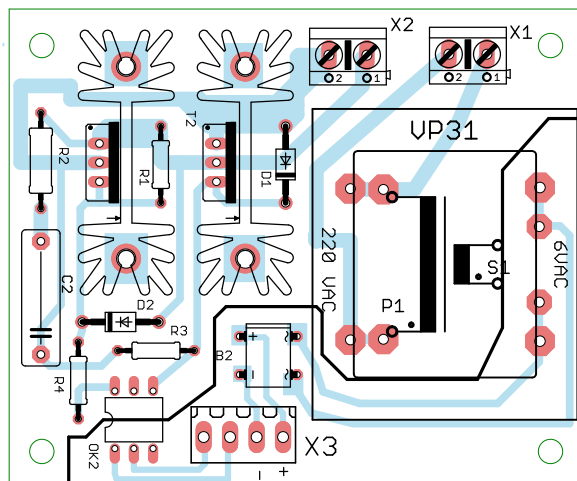
1. Schemat ideowy sterownika (płytki wykonawcza)



1. Schemat ideowy sterownika (płytki sterownika)

zabezpieczającą D2, podający napięcie pulsujące na wejście mikrokontrolera; w dalszej kolejności napięcie jest filtrowane i stabilizowane do 5 V trójkońcówkowym stabilizatorem IC1 (7805).

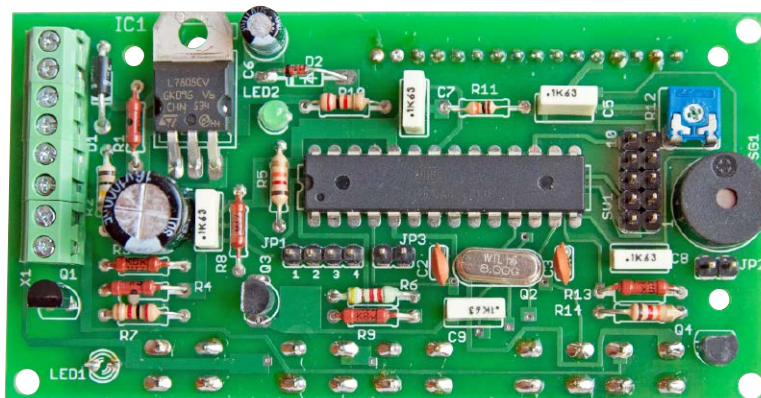
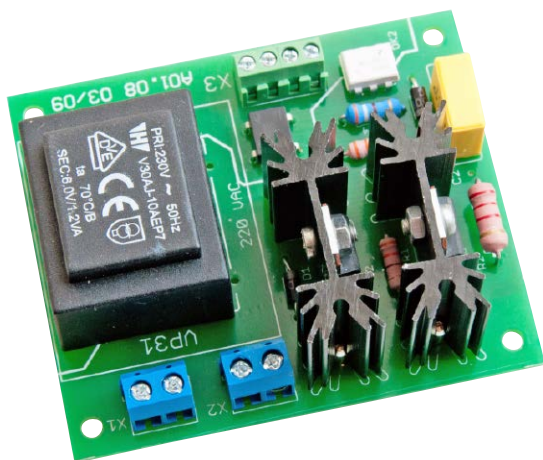
Mikrokontroler IC2 (ATmega8), taktowany rezonatorem kwarcowym, współpracuje z pięcioma przyciskami i wyświetlaczem LCD 2×16. Dwa przyciski zmieniają czas zgrzewania, dwa kolejne – opóźnienie zadziałania, a ostatni



2. Rozmieszczenie elementów na płytkach (wykonawczej i sterownika)

Tabela 1. Wykaz elementów zestawu AVT5553

Grupa	Elementy
Rezystory	R1, R4, R9, R13: 4,7 kΩ; R2: 100 kΩ; R3, R5, R8, R11, R14: 10 kΩ; R6: 220 Ω; R7, R10: 2,2 kΩ; R12: 56 Ω; R15: 390 Ω/0,5 W; R16, R17: 330 Ω/1 W; R18: 2,2 kΩ/2 W; PR1: potencjometr montażowy 10 kΩ
Kondensatory	C1: 1000 μF; C2, C3: 27 pF; C4, C5, C7–C9: 100 nF; C6: 100 μF; C10: 100 nF/250 VAC
Półprzewodniki	IC1: 7805; IC2: ATmega8; T1, T3: BC337; T2: BC327; D1, D3, D4: 1N4007; D2: 1N4148; LED1: czerwona 3 mm; LED2: zielona 3 mm; OK1: MOC3021; B1: mostek 0,5 A/50 V; TY1, TY2: TYN616 (+ radiator); DIS1: LCD 2×16
Inne	S1–S5: przycisk 13,5 mm; SG1: buzzer z generatorem 5 V; TS1: transformator 2–2,5 VA/6 VAC; Q1: rezonator 8 MHz; X1, X2: DG301-5.0/3; X3, X4: DG360-7.5/2; pozostałe elementy montażowe
Opcjonalne	termistor NTC 10 kΩ; wentylator 5 V DC



Zmontowane płytki: sterownika i wykonawcza

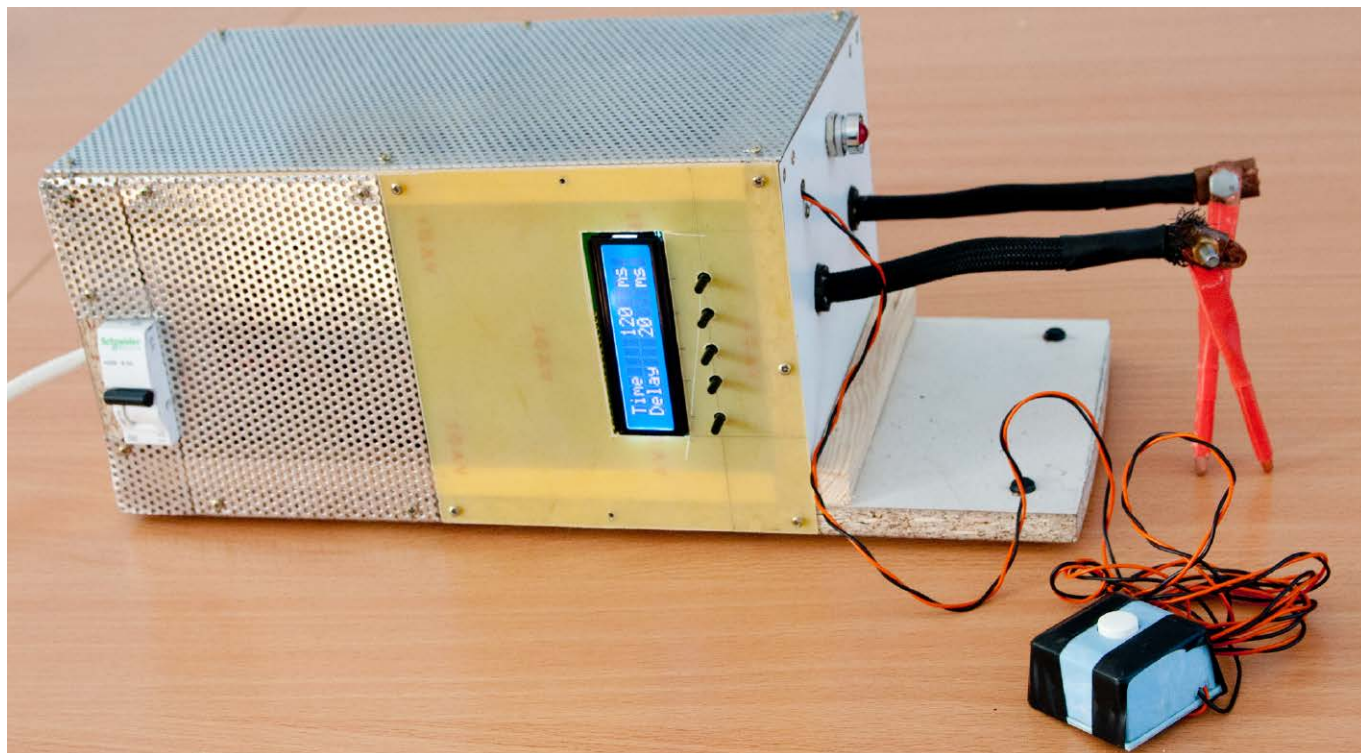
załącza transformator. Obwody dodatkowe: tranzystor T3 steruje brzęczykiem sygnalizującym zgrzewanie, T2 załącza diodę LED w optoizolatorze, a T1 włącza opcjonalny wentylator. Sygnał sterujący wentylatorem dostarcza termistor NTC 10 kΩ włączony między pin 4 złącza X1 a masę.

Na płycie wykonawczej umieszczono transformator sieciowy z mostkiem, zasilający sterownik, oraz przełącznik tyrystorowy na dwóch tyrystorach (TY1, TY2 – TYN616) i optotriaku (OK1 – MOC3021) załączanym w dowolnym momencie. Kondensator C10 i rezystor R18 tłumią przepięcia, które mogłyby uszkodzić tyrystory.

Montaż i uruchomienie

Rozmieszczenie elementów pokazano na rysunku 2. Montaż jest typowy i nie wymaga szczególnego omówienia. Płytkę sterownika ma elementy po obu stronach laminatu: na jednej wyświetlacz LCD, przyciski sterujące i diody LED, na drugiej pozostałe elementy. Wyświetlacz lutuje się do płytki przez listwę kołkową typu goldpin.

Uruchomienie sterownika sprowadza się do włączenia zasilania i ustawienia kontrastu potencjometrem PR1. Po naciśnięciu przycisku wyzwolenia powinna zaświecić



się czerwona dioda LED i – przy zwartej zworze JP3 – być słyszalny brzęczyk. Jeśli sterownik zasilimy do testów napięciem stałym, nie ustali momentu przejścia przez zero i odmówi działania, wypisując komunikat błędu; założenie jumpera na piny 1–2 w JP1 pozwala ominąć ten komunikat i wymusić pracę „na ślepo”.

Czas zadziałania reguluje się od 20 ms ze skokiem 20 ms (parzysta liczba okresów sieci podawana na transformator), a opóźnienie – od zera co 10 ms; czasy od góry są praktycznie nieograniczone. Próg załączenia wentylatora ustawiono wstępnie na ok. 40°C; można go zmienić, wciskając jednocześnie dwa przyciski czasu przy załączaniu zasilania i umieszczając termistor w otoczeniu o temperaturze progowej.

Płytkę wykonawczą najlepiej uruchamiać bez transformatora zgrzewarki – w jego miejsce wystarczy włączyć szeregowo zwykłą żarówkę 40...100 W, która powinna zapalać się na zadany czas po naciśnięciu „FIRE”. Wykorzystanie sterownika do innych zastosowań jest możliwe, ale wymaga rozważnego doboru tyrystorów i obwodu tłumiącego; to samo dotyczy zgrzewarek większej mocy – łączenia dwóch, a nawet czterech transformatorów nie testowano.

Bezpieczeństwo. Zgrzewarka jest zasilana z sieci, a płynące prądy nie są małe. Sterownik zasilono napięciem odseparowanym galwanicznie z osobnego transformatora, a konstrukcję podzielono na dwie płytki – przy czym tyrystory i znaczna część płytki wykonawczej są połączone bezpośrednio z siecią, o czym należy pamiętać podczas montażu i prób.

Część B. Komentarz, dyskusja rozwiązań i projekty pokrewne

Część B prowadzimy w trzech wątkach. Najpierw poddamy projekt przykładowy komentarzowi redakcyjnemu, następnie omawiamy spektrum rozwiązań stosowanych w zgrzewarkach oporowych, a na koniec zestawiamy projekty pokrewne, literaturę oraz ofertę kitów AVT.

Komentarz redakcyjny do projektu AVT5553

To rzadki przypadek projektu, który rozprawia się z popularnym mitem – i ma rację. Jego mocną stroną jest świadome, poprawne podejście do dwóch trudnych zagadnień: momentu załączenia i elementu wykonawczego. Poniżej rozwijamy te wątki oraz wskazujemy pole do modernizacji.

Trafny wybór momentu załączenia – dla obciążenia indukcyjnego prawidłowe jest załączanie w maksimum napięcia, a nie w zerze. Powód: przy załączeniu w zerze brakuje siły przeciwelektromotorycznej, więc strumień magnetyczny narasta dwukrotnie powyżej wartości nominalnej, rdzeń się nasyci, indukcyjność pierwotna gwałtownie maleje, a prąd rozruchowy sięga kilkudziesięciu amperów. Załączenie w szczycie napięcia utrzymuje strumień i prąd w granicach normalnej pracy. Sterownik robi to odwrotnie niż popularne układy z optotriakiem detekcji zera (np. MOC3041).

Dwa tyrystory zamiast triaka – para antyrównoległych tyrystorów najlepiej znosi sterowanie obciążeniem

Tabela 2. Porównanie podejść do zgrzewania oporowego

Typ zgrzewarki	Źródło energii	Element wykonawczy	Uwagi
Transformatorowa (AC)	sieć 230 V + transformator	2 tyrystory antyrównoległe/triak	załączać w maksimum; ryzyko udaru i nasycenia
Akumulatorowa (DC)	akumulator 12 V/LiPo	bank MOSFET-ów	sterowanie czasem impulsu w ms, popularne DIY
Kondensatorowa/energetyczna	kondensatory/akumulator	MOSFET z pomiarem prądu	dawkowanie energii, duża powtarzalność

indukcyjnym i jest odporna na zakłócenia wyzwalania; nie ulega też typowym uszkodzeniom triaka (uszkodzona „połowa” zamienia go w prostownik zwierający transformator). Usunięcie mostka diodowego dodatkowo ogranicza oscylacje.

Sterowanie czasem, nie energią – zgrzew jest tu definiowany liczbą okresów sieci (pętla otwarta). Nowoczesne zgrzewarki niskonapięciowe (np. kWeld) mierzą prąd w czasie rzeczywistym i dawkują energię, co daje powtarzalność niezależną od docisku i czystości styku. To naturalny kierunek modernizacji.

Mikrokontroler i interfejs – ATmega8 jest dziś archaiczny; współczesny odpowiednik zyskałby na enkoderze obrotowym, profilach zgrzewu i wygodniejszym menu. Projekt przewiduje już wyzwalanie zewnętrznym przyciskiem (np. nożnym) – wygodne, gdy obie ręce są zajęte.

Jak zgrzewa się oporowo – przegląd rozwiązań

Zgrzewarki transformatorowe (sieciowe AC).

Rozwiązanie z projektu przykładowego: przerobiony transformator z mikrofalówki, sterowany liczbą

okresów sieci. Tanie i mocne, ale kluczowe są moment załączenia (w maksimum dla obciążenia indukcyjnego) i odporny element wykonawczy.

Zgrzewarki akumulatorowe z kluczem MOSFET. Prąd zgrzewania czerpie się z akumulatora samochodowego 12 V lub pakietu LiPo, a załącza go bank równoległych tranzystorów MOSFET sterowany mikrokontrolerem. Czas impulsu nastawia się w milisekundach; to najpopularniejsza dziś szkoła konstrukcji DIY (Malectrics, projekty na Arduino).

Zgrzewarki kondensatorowe/energetyczne. Energię gromadzi bateria kondensatorów lub akumulator, a sterownik dawkuje ją, mierząc prąd w czasie rzeczywistym (np. kWeld). Daje to bardzo powtarzalne zgrzewy, mniej zależne od warunków styku, kosztem bardziej złożonej elektroniki.

Element wykonawczy i moment załączenia. Dla obciążenia indukcyjnego (transformator AC) sprawdzają się dwa tyrystory antyrównoległe i załączanie w maksimum; w układach niskonapięciowych DC królują równoległe MOSFET-y. Metody zestawia tabela 2.

WM

Projekty pokrewne i literatura

Poniższe zestawienie dobrano pod kątem czytelnika zaawansowanego – obejmuje materiały o istotnej wartości warsztatowej i konstrukcyjnej; pominięto wątki forów i mediów społecznościowych.

Zgrzewarki niskonapięciowe (akumulator + MOSFET)

- **Malectrics** – DIY Arduino Battery Spot Welder (kit + projekt open source). sterownik na Arduino Nano, podwójny impuls, bank MOSFET-ów, zasilanie z akumulatora 12 V; nastawa czasu w ms. <https://malectrics.eu/>; kod: https://github.com/KaepnBalu/Arduino_Spot_Welder_V4
- **keenlab** – kWeld – energetyczna zgrzewarka z pomiarem prądu. dawkowanie energii w czasie rzeczywistym dla powtarzalnych zgrzewów. <https://www.keenlab.de/index.php/portfolio-item/kweld/>
- **Instructables** – DIY Arduino Battery Spot Welder. krok po kroku: budowa, część aluminiowa, montaż MOSFET-ów, obudowa. <https://www.instructables.com/DIY-Arduino-Battery-Spot-Welder/>

Tło konstrukcyjne

- **electricbike.com** – Introduction to battery pack design and building (Part 3). wyjaśnienie zgrzewarki na MOSFET-ach zasilanej akumulatorem 12 V oraz uwagi o budowie pakietów. <https://www.electricbike.com/introduction-to-battery-pack-design-and-building-part-3/>

Zgrzewarki i sterowniki mocy w ofercie kitów AVT

Projekt przykładowy jest dostępny jako gotowy zestaw AVT. Poniżej zestawiamy ten kit oraz powiązane układy sterowania mocą ze sklepu AVT (sklep.avt.pl); dostępność do potwierdzenia w sklepie.

- **AVT5553 Sterownik zgrzewarki oporowej.** projekt przykładowy z częścią A – załączanie w maksimum napięcia, dwie płytki, ATmega8.
- **AVT1679 Moduł wykonawczy z triakami.** łączenie dużych mocy z pełną separacją galwaniczną między sterowaniem a obwodem wykonawczym.
- **AVT5067 Regulator mocy 12 A/230 V AC (~2,5 kW).** sterowanie mocą odbiorników sieciowych.
- **AVT1860 Wzmocniony regulator mocy odbiorników 230 V AC.** do większych obciążeń.
- **AVT3218 Regulator mocy z wyświetlaczem LCD.** regulacja z odczytem nastaw.
- **AVT2210 Najprostszy regulator mocy 230 V.** prosty układ na początek.

Od transformatora z mikrofalówki po energetyczne klucze MOSFET – zgrzewarka oporowa pozostaje wdzięcznym projektem warsztatowym, w którym o sukcesie decyduje nie moc, lecz panowanie nad jednym ułamkiem sekundy.

Luckfox Pico 86 Panel – moduł do automatyki domowej z AI, montowany w standardowej puszcze elektrycznej

Rynek automatyki domowej nieustannie się rozrasta, głównie dlatego, iż każdy producent sprzętu AGD dodaje funkcje Smart/IoT do swoich produktów. Na rynku nie brakuje też smartżarówek, smartgniazdek czy smarttermo-
 statów. Luckfox Pico 86 Panel to moduł, który teoretycznie pozwala stworzyć zarówno własne rozwiązanie tego typu, jak i interfejs kontroli innych urządzeń i modułów smart. A jak to wygląda w praktyce?

Specyfikacja i pierwsze wrażenia

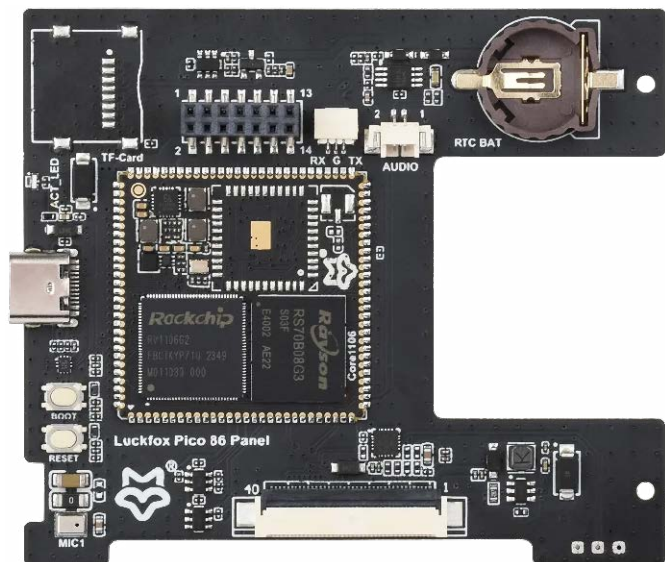
Luckfox Pico 86 Panel to zestaw rozwojowy zaprojektowany z myślą o zastosowaniach smart home i automatyki przemysłowej. Do wyboru są cztery modele: dwa posiadają 128 MB pamięci RAM, a dwa – 256 MB. Dwa modele mają wbudowany moduł Wi-Fi/Bluetooth, a dwa nie. Wszystkie modele mają wbudowaną parę przełączników 230 V/7 A NO, port RS485 i LAN oraz przyłącze zasilania 6...30 V. Moduł można zamontować w puszcze elektrycznej kwadratowej standardu 86 (86×86 mm) o minimalnej głębokości 33 mm. Moduł ma też wbudowany zegar czasu rzeczywistego z możliwością jego podtrzymania baterią. Serce tego panelu jest moduł Luckfox Core1106 z układem SoC Rockchip RV1106 G2 lub G3. Modele 0208 i 1208 zawierają wariant G2, 128 MB pamięci RAM i oferują wydajność w zastosowaniach AI na poziomie 0,5 TOPS. Modele 0408 oraz 1408 używają układu RV1106 G3, mają 256 MB pamięci RAM i oferują większą wydajność AI – 1 TOPS. Modele zaczynające się od cyfry 1 zawierają moduł Wi-Fi/Bluetooth 5.2. Wszystkie modele są taktowane zegarem 1,2 GHz. Wnętrze modułu ukazują **fotografia 1**.

Z zewnątrz moduł sprawia całkiem pozytywne wrażenie. Na froncie znajduje się kwadratowy ekran dotykowy IPS o wymiarach 84×84 mm i rozdzielczości 720×720 px,



z prawej strony znajduje się port USB-C, z lewej zaś kratka głośnika i mikrofonu. Ten panel można odpiąć od tylnej części z terminalami i gniazdem LAN (**fotografia 2**). Panel łączy się z tylną częścią zestawu dzięki prostemu złączu płytka-płytki. Całość wygląda estetycznie, a ekran oferuje szerokie kąty widzenia i przyzwoitą czytelność wewnątrz pomieszczeń. Domyślnie zainstalowana aplikacja demonstracyjna pozwala na połączenie się z siecią Wi-Fi, pokazuje stan połączenia z siecią Wi-Fi i LAN (wraz z przypisanymi adresami IP), podaje datę i godzinę, a także na drugim ekranie oferuje prosty odtwarzacz muzyki z dwiema piosenkami testowymi oraz przyciski do przełączania wbudowanych przełączników. Aplikacja jednak wyświetla niewłaściwą godzinę, prawdopodobnie ze względu na zapisanie danych strefy czasowej w kodzie źródłowym tej aplikacji.

W celu skonfigurowania Wi-Fi jednakże nie użyto aplikacji, lecz podłączono moduł do routera sieci lokalnej przez gniazdo LAN. Używając SSH, nawiązano połączenie



Fotografia 1. Wnętrze modułu

z systemem operacyjnym Buildroot pracującym na module, a następnie ręcznie wprowadzono nazwę sieci Wi-Fi i hasło do pliku konfiguracyjnego. Buildroot to okrojona, minimalistyczna wersja systemu Linux przeznaczona do zastosowań wbudowanych. Ze względu na ograniczenia sprzętowe aplikacje dla Luckfox Pico 86 Panel (i dla całej rodziny modułów Luckfox opartych na układzie RV1106) tworzone są w języku C/C++. Producent zdawał sobie doskonale sprawę z faktu, iż popularny w wielu podobnych platformach Python, nawet w wariacie okrojonym, nie oferuje odpowiedniej wydajności. Za realizację graficznego interfejsu użytkownika odpowiada biblioteka LVGL (*Light and Versatile Graphics Library*). Biblioteka ta jest lekka i zoptymalizowana do współpracy z ekranami dotykowymi. Alternatywą jest również lekka biblioteka MiniGUI. Warto zaznaczyć, iż aplikacja demonstracyjna wygląda i prezentuje się dobrze, mimo swojej prostoty. Szkoda jedynie, iż nie ma większej liczby gotowych aplikacji do zademonstrowania możliwości modułu.



Fotografia 2. Tył modułu ze złączami

Praca z Luckfox Pico 86 Panel

Narzędzia SDK są dostępne dla Ubuntu LTS w wersji 22.04. W systemie Windows SDK można wykorzystać z pomocą narzędzia WSL (*Windows Subsystem for Linux*) oraz trzech rozszerzeń do IDE VS Code. Wymaga to zmiany ustawień na poziomie BIOS-u (uruchomienia funkcji SVM/wirtualizacji dla procesora – funkcja ta czasem jest dość głęboko ukryta), a także instalacji składników systemowych związanych z wirtualizacją. Kolejnym krokiem jest dodanie samego podsystemu Windows dla Linuksa i instalacja odpowiedniej wersji Ubuntu. Po uruchomieniu wiersza poleceń jako administrator należy wpisać komendę:

```
wsl --install --no-distribution
```

Następnie zrestartować system i ponownie włączyć wiersz poleceń jako administrator. Kolejne polecenie to:

```
wsl --install -d Ubuntu-22.04
```

Opcjonalnie można stworzyć dedykowany folder na innej partycji niż systemowa i przenieść tam stworzony właśnie dysk wirtualny poleceniem:

```
wsl --manage Ubuntu --move D:\WSL
```

Rozwiązanie to jest użyteczne, gdy nie chcemy zajmować miejsca na partycji systemowej. Ostatnim krokiem po stronie wiersza poleceń jest aktywacja WSL i wyodrębnienie potrzebnych plików:

```
wsl
git clone https://github.com/LuckfoxTech/luckfox-pico.git
```

Następnym etapem jest konfiguracja środowiska VS Code. Należy zainstalować następujące rozszerzenia:

- WSL;
- C/C++;
- CMake Tools.

W lewym dolnym rogu VS Code znajduje się przycisk kryjący opcję nawiązania połączenia z WSL – po jego nawiązaniu proces kompilacji będzie odbywał się z użyciem pobranego wcześniej i wyodrębnionego zbioru z GitHuba. Zbiór ten zawiera też kilka przykładowych aplikacji demonstracyjnych w formie źródeł. We własnym projekcie należy stworzyć plik toolchain.cmake, którego zawartość pokazuje **listing 1**. Pod linkiem [1] zaś znajduje się kod źródłowy aplikacji demonstracyjnej pracującej na urządzeniu. Po jego przejrzeniu znaleziono przyczynę nieprawidłowego wyświetlania

```

set(CMAKE_SYSTEM_NAME Linux)
set(CMAKE_SYSTEM_PROCESSOR arm)

# Ścieżka do pobranego toolchaina z SDK Luckfox
set(TOOLCHAIN_DIR „/home/<nazwa użytkownika>/luckfox-pico/tools/linux/toolchain/arm-rockchip830-linux-
uclibcgnueabihf”)

set(CMAKE_C_COMPILER „${TOOLCHAIN_DIR}/bin/arm-rockchip830-linux-uclibcgnueabihf-gcc”)
set(CMAKE_CXX_COMPILER „${TOOLCHAIN_DIR}/bin/arm-rockchip830-linux-uclibcgnueabihf-g++”)

```

Listing 1. Zawartość toolchain.cmake

```

void custom_init()
{
    // Time Zone
    setenv(„TZ”, „UTC-8”, 1);
    tzset();

    // Relay
    gpio_export(32);
    gpio_out_direction(32);

    gpio_export(33);
    gpio_out_direction(33);

    // Music
    system(„mpv 2>&1 >/dev/null”);
    music_player_thread_init();
}

```

Listing 2. Kluczowy fragment kodu źródłowego UI_Custom.c odpowiadająca za błędne wyświetlanie godziny na panelu

```

#!/bin/sh
[ -f /etc/profile.d/RkEnv.sh ] && source /etc/
profile.d/RkEnv.sh
case $1 in
start)
    if [ -f /usr/bin/t_s ]; then
        /usr/bin/t_s &
    elif [ -f /usr/bin/86UI_demo ]; then
        /usr/bin/86UI_demo &
    fi
    ;;
*)
    exit 1
    ;;
esac

```

Listing 3. Zawartość pliku /etc/init.d/S99lvgl

godziny – aplikacja faktycznie narzuca swoją strefę czasową, co widać na listingu 2.

Alternatywnym rozwiązaniem jest użycie Ubuntu w wersji 22.04 jako środowiska programistycznego, na przykład za pomocą maszyny wirtualnej Oracle VM VirtualBox. Nie powinno też być problemów z użyciem narzędzi firmy Luckfox w najnowszej wersji Ubuntu, a w dystrybucjach z innych rodzin zawsze można użyć Dockera lub któregoś z rozwiązań VM.

Jeśli chcemy wykonać własną aplikację na module, trzeba ją wstępnie przesłać do podfolderu /usr/bin/ na urządzeniu, używając albo narzędzia ADB (zwykle stosowanego dla urządzeń z Androidem), albo protokołu SCP. W następnej kolejności należy zedytować plik

/etc/init.d/S99LVGL, w którym znajduje się kod startujący domyślną aplikację 86UI_demo. Zawartość tego pliku przedstawia listing 3.

Producent nie dostarcza ujednoliconej dokumentacji, tylko kilka tutoriali na swojej stronie Wiki oraz wspomniane wcześniej przykłady. Jest to jak najbardziej wystarczające dla średnio zaawansowanych programistów, szczególnie obeznanych zarówno ze standardowymi wersjami Linuksa, jak i przeznaczoną dla systemów embedded dystrybucją Buildroot. Przydałaby się aplikacja testowa/demo używająca wbudowanego w układ RV1106 rdzenia NPU. Modele tworzone są z użyciem Pythona, a potem konwertowane do formatu RKNN. Po stronie modułu dostępne są API zarówno dla C, jak i Pythona. Warto nadmienić, iż moduł posiada zainstalowane środowisko Python w wersji 3.11.6 (w posiadanym egzemplarzu).

Podsumowanie

Luckfox Pico 86 Panel to ciekawa propozycja dla chcących produkować systemy typu smart home czy unowocześniać systemy/maszyny przemysłowe bez konieczności projektowania rozwiązania sprzętowego od zera. Jednakże dla kogoś chcącego dodać odrobinę sztucznej inteligencji albo ładnie wyglądający panel kontrolny do własnego systemu automatyki domowej to rozwiązanie może być nieco zbyt skomplikowane od strony oprogramowania. Jest tylko garść przykładów i tylko jedna gotowa aplikacja demonstracyjna zainstalowana już na urządzeniu. Przy ograniczonej liczbie dostępnych zasobów hobbyści będą mieli spore problemy, by wykorzystać ten panel. Szkoda, bo sama konstrukcja jest ładnie zaprojektowana i dobrze przemyślana pod kątem sprzętu i możliwości. Producent nieco zmarnował potencjał urządzenia, nie tworząc większej liczby bardziej rozbudowanych, gotowych aplikacji.

Paweł Kowalczyk

Lista odnośników:

[1] https://github.com/LuckfoxTECH/luckfox_pico_86panel_ui_demo

**Najważniejsze parametry:**

- monofoniczny wzmacniacz audio pracujący w klasie D
- moc wyjściowa: do 1,4 W przy zasilaniu 5 V i impedancji 8 Ω
- rozmiar płytki tylko 20 mm $\times$ 20 mm
- bez dodatkowego radiatora
- sprawność do 84%
- możliwość sterowania głośnikami o impedancji 4 Ω , 16 Ω i 32 Ω
- pobór prądu w stanie spoczynku do 4,5 mA (zasilanie 5,5 V)
- pobór prądu przy pełnymysterowaniu do 270 mA (zasilanie 5 V)
- zasilanie napięciem stałym 2,5...5 V (głośnik 8 Ω) lub 2,5...4,2 V (głośnik 4 Ω)

* **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 wymagają 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):

- Stereofoniczny wzmacniacz klasy D o mocy 2 $\times$ 50 W
- Wzmacniacz z kanałem basowym 2:1
- Miniaturowy wzmacniacz mocy z LM4952
- Miniaturowy wzmacniacz z układem TDA7233S
- Miniaturowy wzmacniacz słuchawkowy
- Cyfrowy wzmacniacz audio 2 $\times$ 10 W w formacie RPi Zero
- Bateria mini wzmacniacz mocy
- Miniaturowy wzmacniacz o mocy 2 $\times$ 6 W

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

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

W ofercie AVT*

AVT6101

Miniaturowy wzmacniacz audio w klasie D

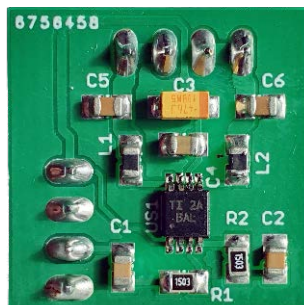
Czy na powierzchni wielkości opuszka kciuka da się zmieścić kompletny wzmacniacz mocy? 70 lat temu to pytanie zostałooby zbyt niedowierzającym spojrzeniem. Tymczasem w XXI wieku mamy technologię, która pozwala taki kompletny moduł wykonać.

Kompletny moduł wzmacniacza audio na jednostronnej płytce drukowanej o wymiarach 20 mm $\times$ 20 mm, bez jakiegokolwiek dodatkowego chłodzenia? To nie żart, lecz fakt. Dzięki temu małemu i lekkiemu modułowi można bez problemu „udźwiękować” takie rzeczy jak sprzęt AGD, zabawki czy domofony. Niemal nie wydziela ciepła, więc nie trzeba szczególnie istotnie przejmować się chłodzeniem go. A – co najważniejsze – na „pokładzie” ma wszystkie niezbędne elementy, nie trzeba dołączać żadnych filtrów czy kondensatorów elektrolitycznych

Budowa

Schemat ideowy omawianego układu znajduje się na rysunku 1. Głównym podzespołem jest układ TPA2005, będący gotowym modułem wzmacniacza w klasie D. Na rysunku 2 znajduje się jego uproszczony schemat blokowy. Zawiera w sobie wszystko to, czego potrzebuje wzmacniacz pracujący impulsowo: przedwzmacniacz, generator PWM, wyjściowy mostek H oraz zabezpieczenia przez zwarcie i przegrzaniem. Elementy zewnętrzne służą ustaleniu wzmocnienia oraz odsprzęgnięciu zasilania.

Rezystory R1 i R2 ustalają wzmocnienie napięciowe układu na 1 V/V. Posiada on wejście różnicowe, stąd na jedno wejście trafia sygnał, a drugie jest „sygnałowo” zwarte z masą, stąd wzmocnienie wynikające z obliczeń wynosi 2 V/V, ale tylko jedno wejście jest sterowane, dlatego układ „widzi” sygnał wejściowy o dwukrotnie mniejszej amplitudzie. Kondensatory C1 i C2 odcinają składową stałą. Można zwiększyć wzmocnienie



układu przez zmniejszenie rezystancji R1 i R2 do np. 15 k Ω , co zwiększyłoby wzmocnienie do 10 V/V, choć polecam przy tym zwiększenie pojemności C1 i C2 do 1 μ F dla zachowania pasma przenoszenia od dołu.

Wejście SHUTDOWN, w które jest wyposażony niniejszy układ scalony, nie jest w tym projekcie wykorzystywane, dlatego spolaryzowano je potencjałem dodatniej linii

Wykaz elementów:**Rezystory:**R1, R2: 150 k Ω 1% SMD0805 (opis w tekście)**Kondensatory:**

C1, C2: 100 nF SMD0805 (opis w tekście)

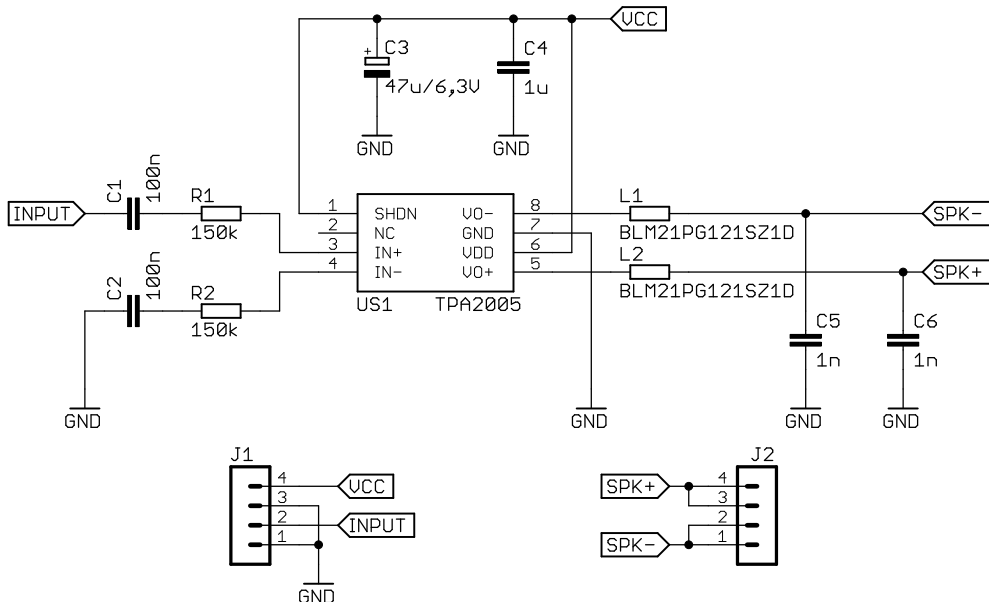
C3: 47 μ F 6,3 V tantalowy SMD AC4: 1 μ F SMD0805**Półprzewodniki:**

U1: TPA2005D1DGN

Pozostałe:

J1, J2: goldpin 2 pin męski prosty 2,54 mm THT (opis w tekście)

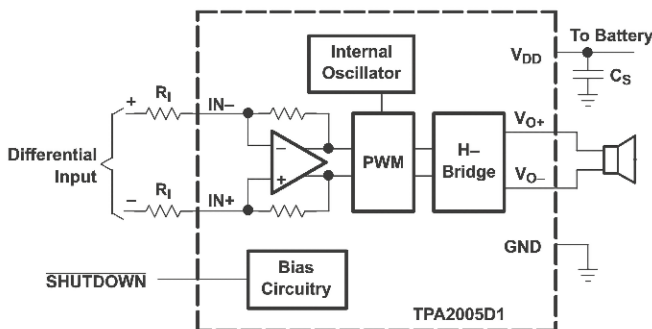
L1, L2: BLM21PG121S21D lub podobny SMD0805



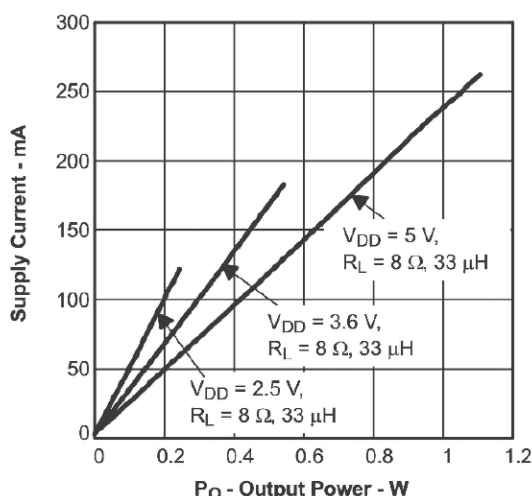
Rysunek 1. Schemat ideowy miniaturowego wzmacniacza audio

zasilającej. Do odsprężania użyto dwóch kondensatorów: ceramicznego $1\ \mu\text{F}$ i tantalowego $47\ \mu\text{F}$. Użycie kondensatorów o większej pojemności mija się z celem, ponieważ częstotliwość kluczkowania wynosi $250\ \text{kHz}$, a to właśnie przełączające się tranzystory wyjściowe są głównym źródłem tętnień na linii zasilania.

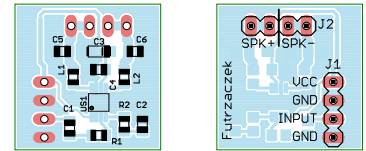
Producent reklamuje ten układ jako „filter-free”, co oznacza, że nie ma konieczności stosowania filtrów LC na wyjściu,



Rysunek 2. Schemat blokowy układu TPA2005 [1]



Rysunek 4. Zależność między poborem prądu a oddawaną mocą [1]



Rysunek 3. Schemat montażowy płytki

których zadaniem byłoby stłumienie sygnału wielkiej częstotliwości i jego harmonicznych, wynikających wprost z kluczowania. Jednak w nocie katalogowej [1] można znaleźć adnotację, że warto zastosować filtr wyjściowy, o ile przewody łączące układ z głośnikiem są krótkie – poniżej 100 mm przeszedł

testy kompatybilności elektromagnetycznej. Dlatego doda-
łem na płytce prosty filtr, dopuszczalny przez producenta,
składający się z koralików ferrytowych i kondensatorów ce-
ramicznych o pojemności 1 nF. Taki zabieg przysłuży się
ograniczeniu emitowanych przez układ zakłóceń, a koszt
takiego rozwiązania jest wręcz symboliczny.

Montaż i uruchomienie

Układ został zmontowany na jednostronnej płytce drukowanej o wymiarach 20 mm × 20 mm. Jej wzór ścieżek oraz schemat montażowy przedstawia **rysunek 3**. Wszystkie elementy SMD są montowane na stronie dolnej, na górnej wystają jedynie złącza szpilkowe typu goldpin, użyte w prototypie. Nic nie stoi na przeszkodzie, by pominąć złącza J1 i J2, lutując przewody bezpośrednio w płytkę, co jeszcze bardziej zmniejszy wysokość tego modułu.

Prawidłowo zmontowany układ jest gotowy do działania od razu po podłączeniu głośnika (do wyprowadzeń SPK+ i SPK–, zdublowane dla wygody), sygnału wejściowego (do zacisków INPUT i GND) oraz zasilania (do zacisków GND i VCC). Układ TPA2005 najlepiej radzi sobie z obciążeniem głośnikiem o impedancji nominalnej 8 Ω , ale przy obniżonym napięciu zasilania (do 4,2 V) może działać z głośnikami o impedancji 4 Ω . Wysterowanie odbiorników o większej impedancji, jak 16 Ω czy 32 Ω , nie stanowi dla niego problemu, trzeba jedynie liczyć się z niższą mocą wyjściową z racji ograniczonego napięcia zasilania (do 5,5 V). Pobór prądu zależy od aktualnie oddawanej mocy – szczegóły na **rysunku 4**.

Michał Kurzela, EP

Źródła:

[1] <https://www.ti.com/lit/ds/svmlink/tpa2005d1.pdf>

Oscyloskop, który wchodzi w wysokie napięcie

Portfolio pomiarowe Cleverscope i TiePie engineering w ofercie Egmont Instruments

Oscyloskop od dawna nie jest już wyłącznie pudełkiem z ekranem. Rosnąca klasa przyrządów PC-based – bez własnego wyświetlacza, za to z przetwornikiem A/C o rozdzielczości 14–16 bitów i całą „aparaturą” przeniesioną do programu na komputerze – znalazła sobie mocne nisze tam, gdzie klasyczny oscyloskop stołowy sobie nie radzi: w izolowanych pomiarach elektroniki mocy, w wielokanałowej akwizycji i w diagnostyce warsztatowej. W ofercie firmy Egmont Instruments znajdziemy dwa uzupełniające się rody takich przyrządów: nowozelandzki **Cleverscope** i holenderski **TiePie engineering**. Oba mają na polskim rynku długą historię, a rok 2026 przynosi w nich kilka nowości, które warto poznać.

Dwie filozofie jednego przyrządu

Choć obie marki należą do tej samej rodziny „oscylloskopów w laptopie”, reprezentują dwie różne strategię. Cleverscope idzie w głąb jednego, trudnego problemu – izolowanego pomiaru w elektronice mocy dużych napięć, aż po eksperymentalne układy z węgla krzemu (SiC) i azotku galu (GaN). TiePie idzie wszędzie – buduje jedną platformę sprzętowo-programową, w której ten sam przyrząd jest naraz oscyloskopem, analizatorem widma, rejestratorem, multimetrem, generatorem i analizatorem protokołów, i skaluje ją od taniego miernika sieciowego po wielokanałowe systemy akwizycji i diagnostykę samochodową.

Wspólny mianownik jest jednak wyraźny: przyrząd definiowany programowo (software-defined), wysoka

rozdzielczość pionowa, głęboka pamięć, nacisk na izolację galwaniczną i na elektronikę mocy. To właśnie ona – napędzana dziś przez szybkie tranzystory SiC/GaN – wymusza pomiary, których zwykły oscyloskop wykonać nie potrafi: na „gorącej” stronie, przy zmianach napięcia rzędu 100 V/ns i napięciach wspólnych liczonych w kilowoltach.

Cleverscope – specjalista od izolacji i elektroniki mocy

Cleverscope Ltd działa z Auckland w Nowej Zelandii od 2005 roku i od początku stawia na jeden pomysł: oscyloskop mieszany (mixed-signal) klasy PC z 14-bitowym przetwornikiem, izolowanymi kanałami, wbudowanym generatorem i funkcją analizy odpowiedzi częstotliwościowej



Cleverscope CS548 z wpiętym zdalnym modułem CS1200 (IsoPod)

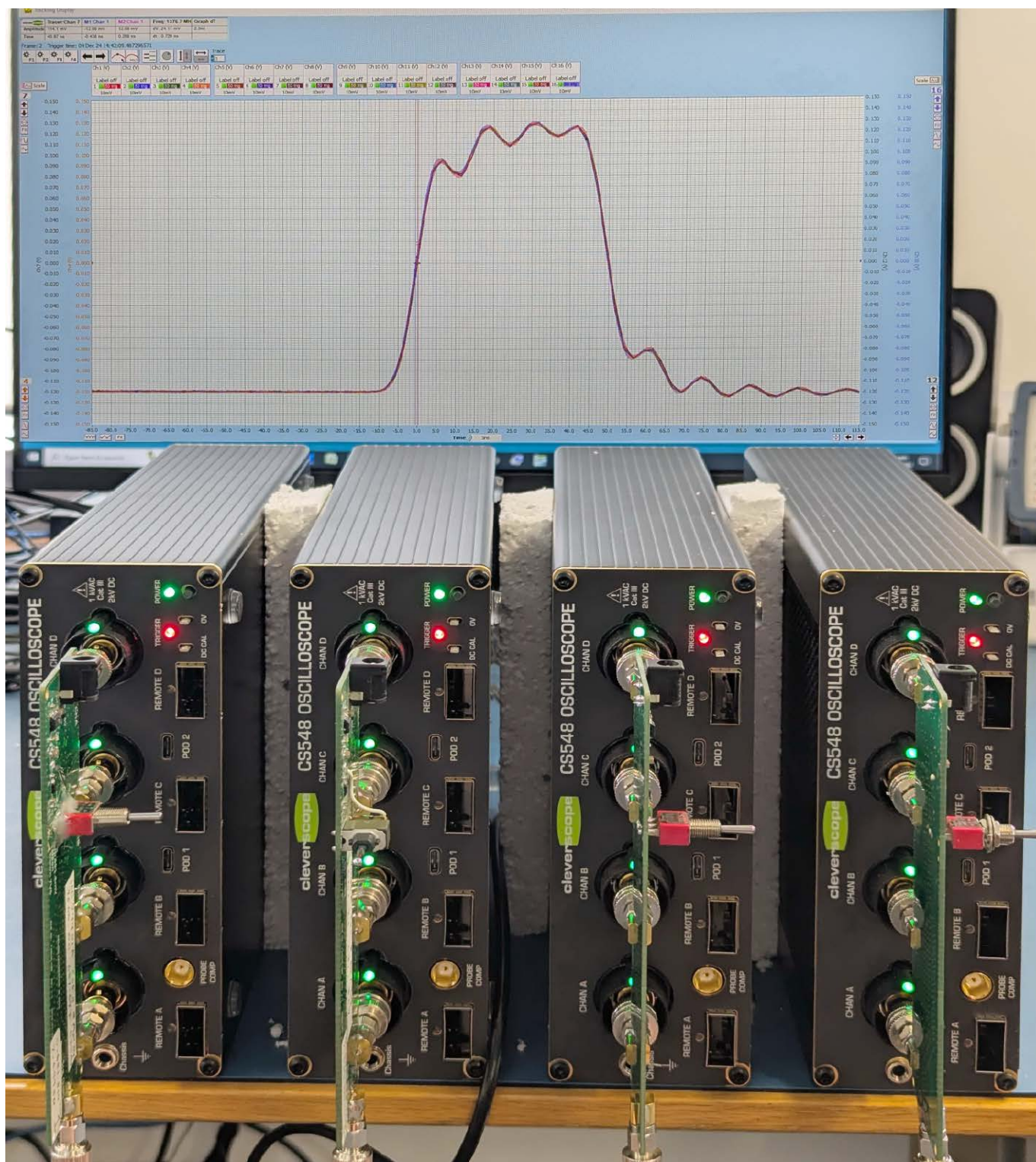
Czym jest FRA?

Analiza odpowiedzi częstotliwościowej (*Frequency Response Analysis*) to pomiar wzmocnienia i przesunięcia fazowego układu w funkcji częstotliwości. W zasilaczach impulsowych pozwala wyznaczyć zapas fazy i wzmocnienia pętli sprzężenia zwrotnego, czyli ocenić stabilność przetwornicy jeszcze na stole laboratoryjnym. Do pomiaru potrzebny jest generator pobudzający pętlę oraz dwa kanały odczytujące odpowiedź – dlatego wbudowany, izolowany generator CS548 jest tu tak istotny.

(FRA). Całością steruje 64-bitowa aplikacja **CS302**, która potrafi połączyć od jednego do czterech przyrządów w jeden oscyloskop o 4...16 kanałach, ze streamingiem na dysk, matematyką symboliczną, integracją z Matlabem, dekodowaniem protokołów i statystyką sygnału.

CS548 – izolowany oscyloskop 2 kV

Sztandarowym modelem jest **CS548** – oscyloskop z kanałami izolowanymi do 2 kV, zaprojektowany do zastosowań wszędzie tam, gdzie potrzebna jest izolacja i wysoki współczynnik tłumienia sygnału wspólnego (CMRR),



a w szczególności w elektronice mocy. Ma cztery wewnętrzne kanały izolowane, znoszące zmianę napięcia wspólnego 100 V/ns przy napięciu wspólnym do 2 kV, oraz rozdzielczość 14 bitów przy szumie na poziomie 1 części na 1300. Kluczem jest wbudowany **izolowany generator 0...65 MHz**, służący do analizy FRA: badania stabilności pętli, charakterystyk wzmocnienia i fazy pracujących modułów mocy, impedancji szyn zasilających oraz pomiaru indukcyjności i pojemności w funkcji częstotliwości.

W komplecie z CS548 dostajemy m.in. dla każdego kanału sondę 1x/10x i sondę 100x (± 800 V, opcjonalnie 200x do ± 1600 V), zaciskowy dławik trybu wspólnego poprawiający CMRR przy dużych szybkościach narastania, dwa izolowane moduły wejść cyfrowych CS1301 (razem 8 wejść) oraz adapter SFP (Ethernet miedziany lub światłowodowy). Cena katalogowa to 13 500 USD. Dla użytkowników, którzy chcą pracować wyłącznie na kanałach zdalnych, dostępna jest odmiana CS508 – bez kanałów wewnętrznych – za 5000 USD.

Kanały zdalne – „sondy”, które są całym torem cyfrowym

To, co w katalogu bywa nazywane „sondami”, to w istocie zdalne moduły akwizycji (producent nazywa je IsoPod). **CS1200** (kanał napięciowy) i **CS1201** (kanał prądowy) to zasilane z baterii moduły łączone z CS548 aktywnym światłowodem w standardzie QSFP na odległość od 1 m do 30 m. Takie łącze daje izolację do 30 kV i CMRR ponad 100 dB w całym paśmie. Sam CS1200 próbkuje z szybkością 500 MSa/s przy paśmie 200 MHz i rozdzielczości 14 bitów, ma szum na poziomie 1 części na 5000 i dokładność 0,15%. Oba moduły kosztują po 3730 USD, a długość światłowodu wybiera się przy zamówieniu.

Ofertę uzupełniają izolowane moduły cyfrowe CS1301, moduł wejść/wyjść CS1302 (generowanie sekwencji double-pulse oraz sterowania PWM w układach półmostka i pełnego mostka) oraz moduł pomiaru napięcia nasycenia CS1133 do pomiaru strat mocy przy przełączaniu do 3 kV.

Nowości 2026 – od „całego tranzystora” po 10 kV

Najciekawsze w tegorocznej ofercie Cleverscope są dwie zapowiedzi, obecnie w fazie przedprodukcyjnej:

- **CS1202 Transistor Digitizer** („Capture a whole transistor”) – moduł zbudowany wokół techniki DSRT (Dynamic Sample Rate Technology). Standardowy interwał próbkowania wynosi 4 ns, ale w oknie odniesionym do momentu wyzwolenia skraca się do 200 ps na czas 408 ns. Dzięki temu minimalizowana jest zajętość pamięci i pobór mocy przetwornika, a sprzętowy mnożnik (PSW)

SST – transformator półprzewodnikowy

Solid State Transformer to transformator, w którym klasyczny rdzeń sieciowy 50 Hz zastępuje się przekształtnikiem energoelektronicznym pracującym na wysokiej częstotliwości. Pozwala to radykalnie zmniejszyć masę i gabaryty oraz aktywnie sterować przepływem energii – stąd zainteresowanie w sieciach inteligentnych, ładowaniu pojazdów i systemach OZE. Ceną jest praca przy napięciach rzędu kilku-kilkunastu kV i bardzo szybkich zboczach, co wymaga aparatury izolowanej na 10 kV i o paśmie licznym w setkach MHz – dokładnie takiej, jaką zapowiada Cleverscope.

i sprzętowy integrator (ESW) pozwalają uchwycić moc szczytową i energię każdego przełączenia przy częstotliwościach do 100 kHz i zmienności wypełnienia 1...99%.

- **CS1203 10 kV Channel** – kanał mierzący wychYLENIA napięcia do 10 kV: pasmo 200 MHz, zakresy ± 1 kV i ± 10 kV, rozdzielczość 125 mV/1,25 V, CMRR ponad 100 dB przy 50 MHz, połączenie przewodem w izolacji silikonowej. To podstawa systemu Cleverscope do badania transformatorów półprzewodnikowych (SST). W przygotowaniu jest jeszcze CS1204 VON, mierzący napięcie tranzystora w stanie przewodzenia przy swingu do 10 kV.

Razem z istniejącymi kanałami CS1200 i CS1201 tworzą one komplet do pomiarów w układach SST: strat przełączania i przewodzenia, rezystancji w stanie włączenia (RDS(on)), indukcyjności szyny i parametrów samego przyrządu półprzewodnikowego.



Idea pomiaru w systemie SST: zapowiadany kanał 10 kV (CS1203) przy przyrządzie SiC

Tabela 1. Wybrane modele z obu rodzin (dane katalogowe producentów)

Model	Kanały	Kluczowe parametry toru	Izolacja/uwagi	Cena netto*
Cleverscope CS548	4 izol. (+4 zdalne)	14 bit; wbud. izol. generator 0...65 MHz (FRA); streaming	CMV 2 kV, dV/dt 100 V/ns	13 500 USD
Cleverscope CS1200 (kanał zdalny)	1 izol.	500 MSa/s, 200 MHz, 14 bit	izolacja 30 kV, CMRR >100 dB, światłowód 1...30 m	3730 USD
Cleverscope CS1203 (zapowiedź)	1 izol.	200 MHz; zakres $\pm 1/\pm 10$ kV; CMRR >100 dB	podstawa systemu SST (10 kV)	wkrótce
TiePie Handyscope HS6 DIFF	4 różn.	1 GSa/s, 250 MHz, 256 Mpts	wejścia różnicowe	od 1434 €
TiePie Handyscope TP450	1 różn. izol.	250 kSa/s, 200 kHz, 16 bit	izol. galw. 5000 Vrms, wejście do 450 V	od 163 €
TiePie Wi-FiScope WS6 DIFF	4 różn.	1 GSa/s, 250 MHz, 256 Mpts	łącze Wi-Fi/LAN/USB	od 2252 €
TiePie MultiScope MS8	8	do 1 GSa/s, 250 MHz	rozbudowa modułowa CMI/WCMI (do 400 m)	od 2939 €
TiePie ATS610004DW (automotive)	4 różn./SE	1 GSa/s, 250 MHz, 256 Mpts	Wi-Fi/LAN/USB, wsparcie ATIS	3175 €

* Ceny producentów, netto: Cleverscope w USD, TiePie w EUR. Ceny detaliczne w PLN oraz warunki dostawy ustala Egmont Instruments.

TiePie engineering – ekosystem sześciu przyrządów w jednym

TiePie engineering z holenderskiego Sneek reprezentuje przeciwną filozofię: maksymalnej uniwersalności. Sercem oferty nie jest pojedynczy model, lecz wspólne oprogramowanie i modułowa architektura, dzięki którym całą rodzinę można skalować i łączyć.

Multi Channel – sześć przyrządów w jednym oknie

Program **Multi Channel** obsługuje wszystkie przyrządy USB i Wi-Fi TiePie i zmienia każdy z nich w sześć instrumentów: oscyloskop, analizator widma, rejestrator (data logger), multimetr, generator arbitralny (AWG) oraz analizator protokołów. Ten ostatni dekoduje magistrale I²C, UART, CAN, J1939, SPI, LIN, DMX512, SENT i FlexRay. Do dyspozycji są gotowe konfiguracje (Quick

Setups), rozbudowane operacje matematyczne w postaci bloków przetwarzania, praca sieciowa w trybie „plug in and measure” oraz interfejs w wielu językach – w tym po polsku. Uruchomiony bez podłączonego sprzętu program działa jako w pełni funkcjonalne demo.

Handyscope – rdzeń rodziny USB

Podstawą oferty są oscyloskopy USB z serii Handyscope. Najwydajniejsze HS6 DIFF i HS6 to przyrządy czterokanałowe o próbkowaniu do 1 GSa/s, paśmie 250 MHz i pamięci do 256 Mpróbk, przy czym HS6 DIFF ma wejścia różnicowe. Niżej plasują się dwukanałowy HS5 z generatorem AWG, HS4 DIFF i HS4, HS3 oraz kompaktowy Handyprobe HP3. Ceny netto zaczynają się od ok. 129 € (HP3) i sięgają ok. 2858 € za HS6 DIFF. Wszystkie łączą wysoka rozdzielczość i dokładność oraz duża głębokość pamięci.

Egmont Instruments

ul. Marszałkowska 136/31, 00-004 Warszawa

tel. 228506205-07, kom. 692501750

tiepie@egmont.com.pl, <http://www.egmont.com.pl/tiepie>

Egmont

WiFiScope WS4, WS5, WS6, WS6 DIFF przystawki oscyloskopowe DSO + generator AWG + tester EMI

- 2 lub 4 wejścia BNC
- wejścia DIFF lub SE
- próbkowanie do 1 GS/s
- streaming do 200 MS/s
- pasmo do 250 MHz
- rozdzielczość 8, 12, 14, 16 bitów
- zakresy napięć +/-200 mV ... +/-80 V
- pamięć do 256 MS/kanał
- interfejs WiFi 802.11, LAN 1Gb, USB 3.0/2.0
- funkcje: oscyloskop cyfrowy DSO, generator sygnałowy / AWG, tester EMI, analizator widma, woltomierz, data logger / rejestrator, analizator protokołów
- praca synchroniczna wielu modułów
- łącznie dostępnych 108 różnych modeli
- funkcje i parametry zależne od konkretnego modelu





Handyscope TP450
wpięty wprost w gniazdo sieciowe – pomiar jakości energii bez dodatkowego dzielnika

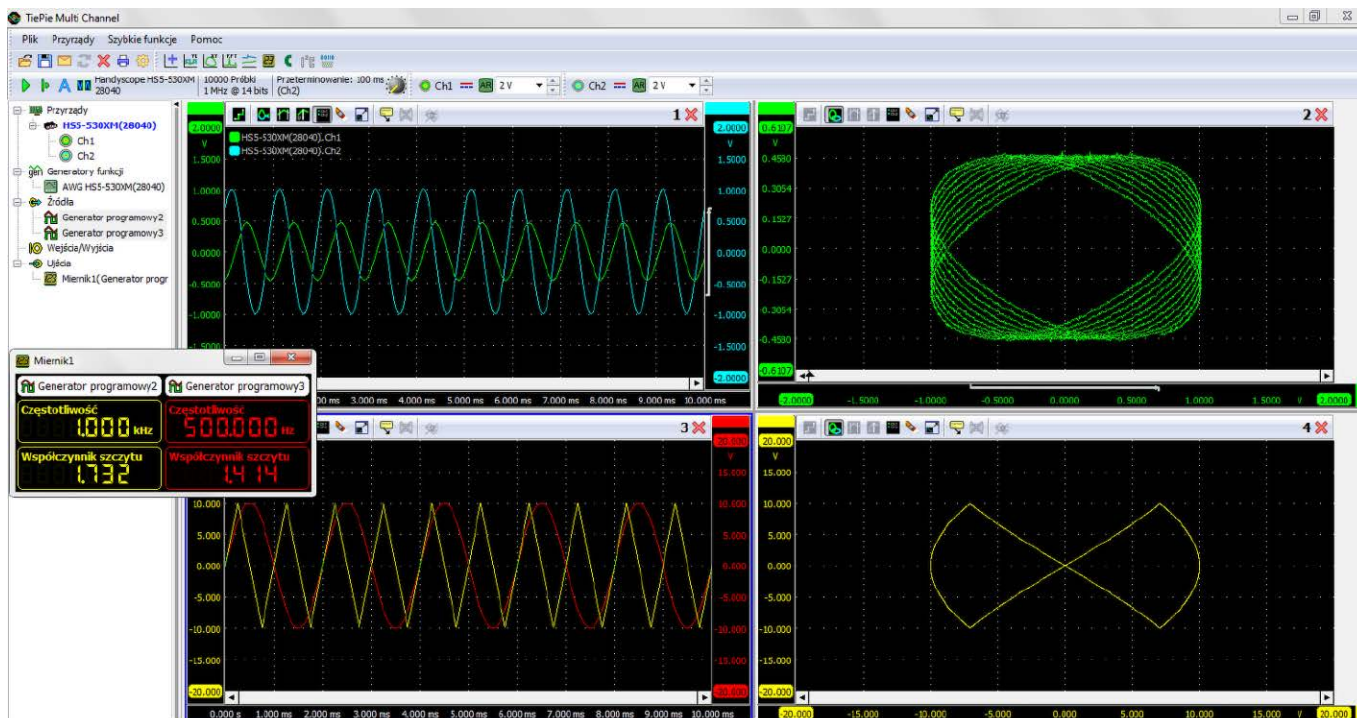
Handyscope TP450 – najtańsze wejście

Osobnej uwagi wart jest **Handyscope TP450** – izolowany galwanicznie analizator jakości sieci za jedyne 163 €. Wtyka się go wprost w gniazdko 110...230 V (zakres wejściowy do 450 V), bez żadnego dodatkowego dzielnika. Ma rozdzielczość 16 bitów (jak woltomierz 5½ cyfry), próbkuje z szybkością do 250 kSa/s – czyli wykonuje 5000 pomiarów na każdy okres sieci – i jest izolowany na 5000 Vrms. Potrafi rejestrować w sposób ciągły przez tygodnie (tydzień to ok. 0,69 GB danych), nie gubiąc ani zapadów, ani

szpilek trwających choćby 1 ms. Trzema takimi modułami zmierzymy sieć trójfazową; poradzi sobie też z pomiarem instalacji 24 V czy wzorców sterowania.

WiFiScope – pomiar na odległość i w izolacji

Rodzina WiFiScope (WS6 DIFF, WS6, WS5, WS4 DIFF) oferuje te same możliwości pomiarowe co Handyscope, ale z łączem sieciowym Wi-Fi/LAN/WAN. Sprawdza się, gdy obiekt jest pod napięciem, w izolacji galwanicznej

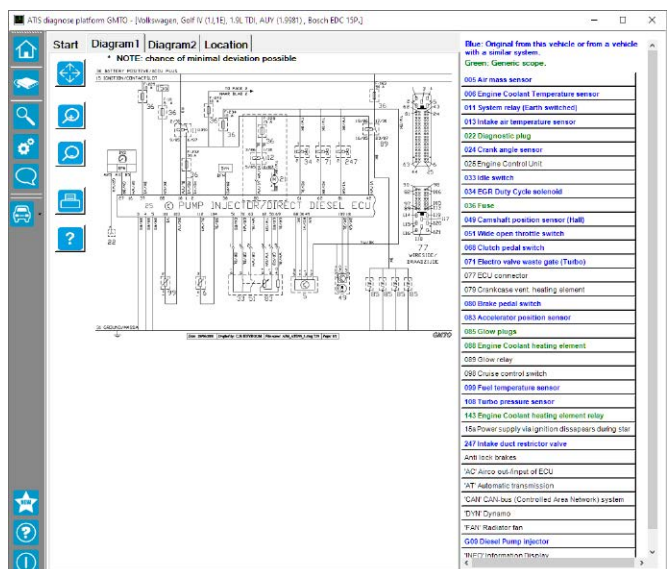
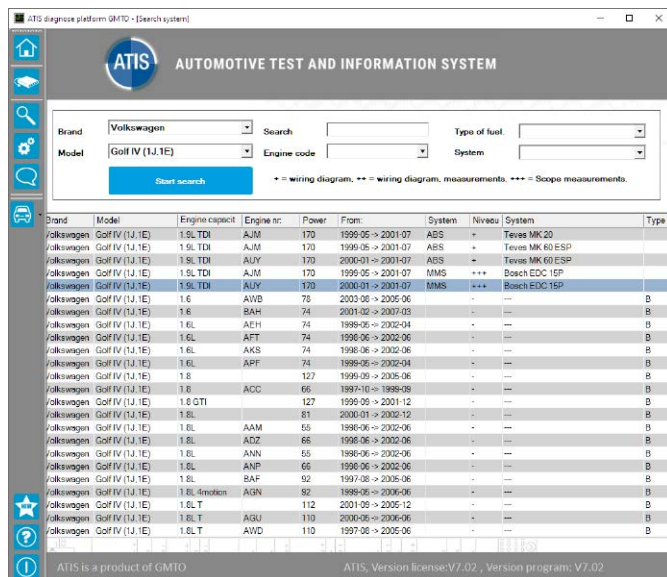


MultiScope: łączenie kilku przyrządów interfejsem CMI w jeden wielokanałowy, w pełni zsynchronizowany system

albo po prostu daleko od stanowiska – program traktuje przyrząd sieciowy tak, jakby był wpięty bezpośrednio do komputera, bez potrzeby specjalistycznej wiedzy o sieci. Ceny mieszczą się w przedziale ok. 2088...2252 €.

MultiScope – własny system wielokanałowy

Gdy cztery kanały to za mało, wkracza **MultiScope** – modułowy zestaw dający od 4 do 12 (a na życzenie i więcej) w pełni zsynchronizowanych kanałów. Powstaje on z kilku przyrządów Handyscope/WiFiScope połączonych interfejsem **CMI** (Combine Multiple Instruments): wystarczy jeden kabel, a układ sam się rozpoznaje, terminuje i synchronizuje z zerowym błędem zegara próbkującego (0 ppm). Wariant bezprzewodowy **WCMI** łączy przyrządy na odległość do 400 m, zachowując izolację galwaniczną. Do budowy własnego oprogramowania i rozwiązań OEM służy zestaw SDK libtiepie-hw. Ceny gotowych konfiguracji zaczynają się od ok. 2516 €.



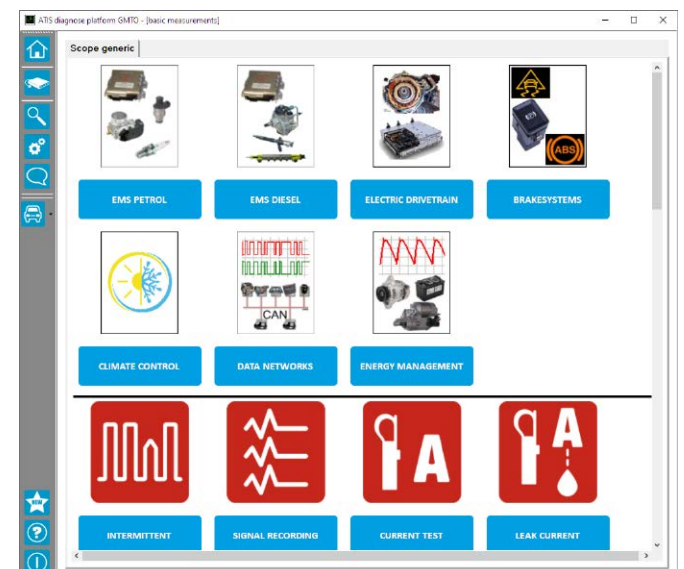
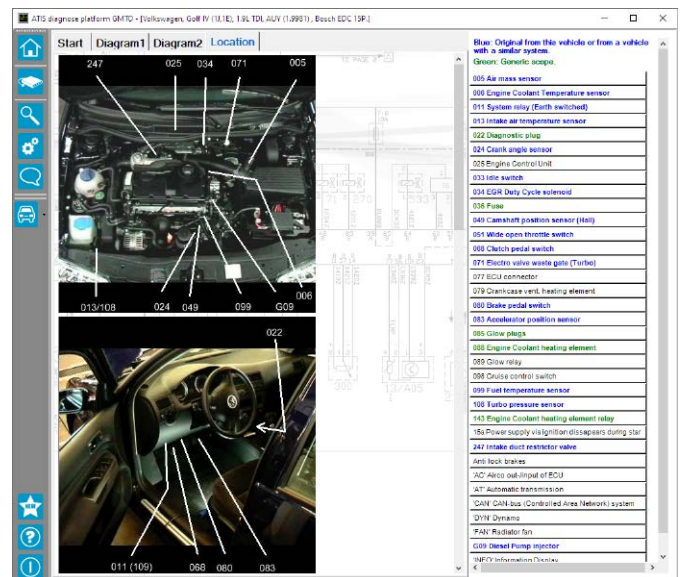
ATIS – baza „znanych-dobrych” przebiegów sprzężona z oscyloskopem diagnostycznym ATS

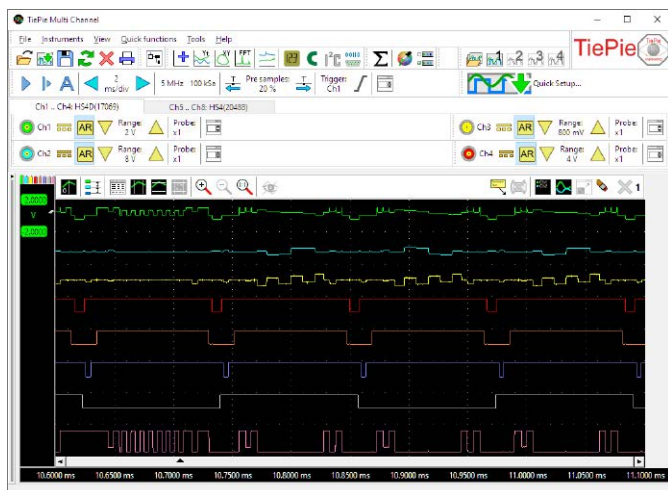
Opcja EMI – wstępne badania kompatybilności

Wybrane modele (m.in. HS6 DIFF) mogą pełnić funkcję testera pre-compliance EMI, czyli pozwalają na wstępną ocenę emisji zaburzeń jeszcze przed wysłaniem wyrobu do akredytowanego laboratorium. Opcja EMI dostępna jest także w innych modelach rodziny, dzięki czemu jeden przyrząd służy zarówno do bieżącej pracy konstruktorskiej, jak i do wczesnego wychwytywania problemów z kompatybilnością elektromagnetyczną.

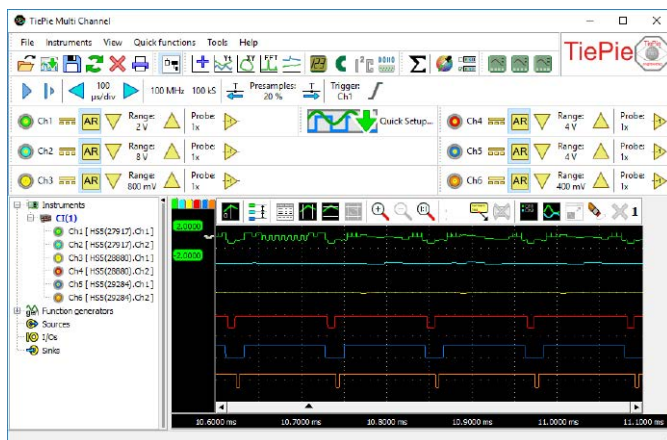
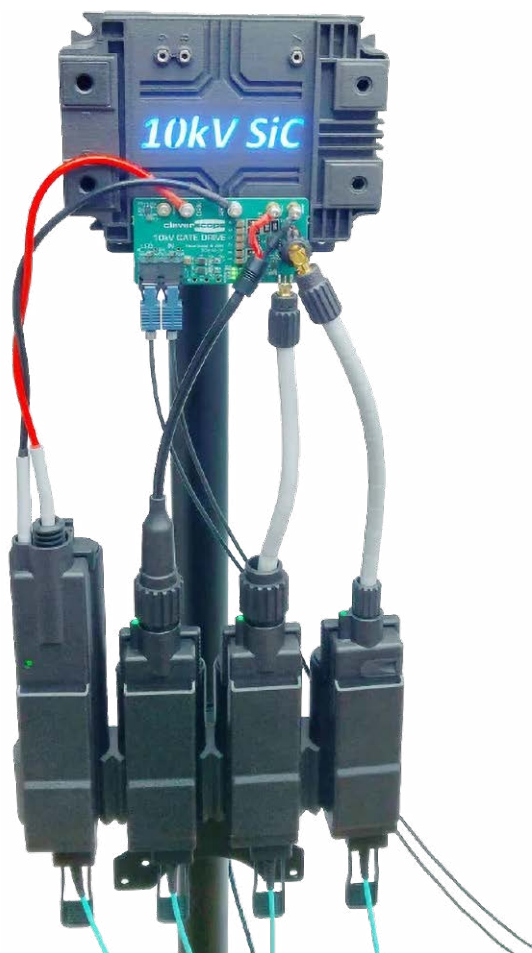
TiePie Automotive i ATIS – oscyloskop w warsztacie

Zupełnie osobną gałęzią jest **TiePie Automotive**. Tworzą ją izolowane oscyloskopy diagnostyczne z serii ATS (Wi-Fi/LAN/USB, cztery kanały różnicowe, próbkowanie do 1 GSa/s, ceny 2093...3175 €) oraz gotowe zestawy diagnostyczne ADK, czyli oscyloskop z akcesoriami, sondami i cęgami prądowymi oraz licencją ATIS Lite. Sercem





systemu jest oprogramowanie **ATIS** (Automotive Test and Information System) – baza wiedzy dla mechanika: schematy, lokalizacje czujników i elementów wykonawczych, opisy oraz przebiegi „znane-dobre” do porównania z własnym pomiarem. Cechą unikalną jest bezpośrednie sprzężenie z oscyloskopem: jedno kliknięcie ładuje właściwe nastawy dla wybranego elementu.



Wersja ATIS Pro (545 €) obejmuje dane dla ponad 3500 pojazdów, z naciskiem na marki europejskie i japońskie, ponad 25 elementami na pojazd, schematami instalacji oraz stronami diagnozy kodów usterek; ATIS Lite (215 €) zawiera informacje generyczne, niezależne od marki. To propozycja nie tyle dla konstruktora, ile dla profesjonalnego warsztatu – pokazuje jednak, jak szeroko można wykorzystać tę samą platformę pomiarową.

Podsumowanie i porównanie

Dwie marki, dwie strategie – a wspólny mianownik pozostaje ten sam: oscyloskop jako element definiowany programowo, o wysokiej rozdzielczości, głębokiej pamięci, z izolacją i mocnym wsparciem dla elektroniki mocy. Cleverscope wybiera specjalizację ekstremalną: 2 kV izolacji w CS548, kanały zdalne na 30 kV, a wkrótce zapowiadane 10 kV do badania transformatorów półprzewodnikowych. TiePie stawia na uniwersalność i ekosystem – od 163 € za analizator sieciowy TP450, przez rodziny Handyscope i WiFiscope, po wielokanałowe MultiScope i kompletną diagnostykę samochodową z bazą ATIS. Wybór zależy więc wprost od zastosowania, a jedno i drugie znajdziemy u tego samego dystrybutora.



Informacje i sprzedaż

Dystrybutor w Polsce: Egmont Instruments

00-004 Warszawa, ul. Marszałkowska 136/31, tel. 22 850 62 05-07; e-mail: egmont@egmont.com.pl; www.egmont.com.pl

Materiały źródłowe: cleverscope.com, tiepie.com, tiepie-automotive.com (stan na lipiec 2026).

Poprzednie dwie części w EP05/2026
i EP06/2026 – dostępne na
www.UlubionyKiosk.pl



Oscyloskopy 2026 – od lampy Brauna do 12 bitów na biurku konstruktora, część 3

W dwóch pierwszych częściach repetytorium prześledziliśmy drogę oscyloskopu od lampy Brauna do współczesnego DSO (rozdziały 1...2) oraz zeszliśmy na poziom próbkowania, wyzwalania, akwizycji i sond (rozdziały 3...4). Pokazaliśmy m.in., dlaczego deklaracja „1 GSa/s w karcie katalogowej” nie zawsze oznacza 1 GSa/s na każdym kanale, kiedy włączyć tryb hi-res, a kiedy peak detect, oraz jak wyzwalanie zaawansowane wyłapuje błąd niewidoczny dla prostego triggera zboczowego.

Trzecia, ostatnia część zamyka cykl i przesuwą akcent z „jak działa oscyloskop” na „jak wybrać właściwy”. Rozdział 5 wyjaśnia, co naprawdę zmieniło się w latach 2020...2026: wejście 12-bitowych przetworników ADC i wbudowanych kanałów logicznych (MSO) do klasy poniżej 10 000 zł. Rozdział 6 przekłada to na praktykę konstruktora – drzewo decyzyjne doboru, szczegółową tabelę

porównawczą reprezentatywnych modeli 2026 r. oraz cztery typowe scenariusze pracy w pracowni.

Punkt wyjścia jest przy tym mniej oczywisty, niż sugeruje marketing. Natywne 12-bitowe DSO nie jest wynalazkiem lat dwudziestych – zadebiutowało już w 2012 r. – ale przez całą minioną dekadę pozostawało w cenie kilkudziesięciu tysięcy złotych. Dopiero lata 2022...2023 zepchnęły

Co znajdziesz w cyklu repetytorium:**Część 1 (EP 05/2026, już opublikowana)**

- Rozdział 1. Krótka historia oscyloskopów
- Rozdział 2. Architektura współczesnego DSO – tor pionowy, ADC i pamięć akwizycji

Część 2 (EP 06/2026, już opublikowana)

- Rozdział 3. Próbkowanie, wyzwalanie i akwizycja – tryby hi-res i pamięć segmentowa
- Rozdział 4. Sondy i integralność sygnału – sondy pasywne i aktywne, pętla masy, kategorie bezpieczeństwa

Część 3 (EP 07–08/2026 – bieżące wydanie)

- Rozdział 5. 12 bitów i MSO w klasie budżetowej – co zmieniło się w latach 2020–2026
- Rozdział 6. Praktyka konstruktora + porównanie modeli reprezentatywnych dla 2026 r.

Archiwalne wydania EP są dostępne na www.ulubionykiosk.pl

natywne 12 bit i MSO do budżetu dostępnego dla pojedynczego konstruktora. Ta część pokazuje, jak do tego doszło, co realnie zyskujemy, a za co płacimy ukrytą cenę w postaci kompromisów konstrukcyjnych.

Materiał opieramy – podobnie jak w częściach 1 i 2 – na dokumentach producenckich i kartach katalogowych: notach Keysighta, Tektroniksa i Rohde & Schwarz, białej księdze Teledyne LeCroy o rozdzielczości oscyloskopów oraz na datasheetach modeli 2020...2026 (Siglent, Rigol, UNI-T, R&S). Bibliografia tej części jest samodzielna i numerowana od [1]; zamieszczamy ją na końcu rozdziału 6.

Rozdział 5. 12 bitów i MSO w klasie budżetowej – co zmieniło się w latach 2020...2026

Na końcu rozdziału 4 zapowiedzieliśmy pytanie, które teraz otwiera trzecią część: jakie kompromisy konstrukcyjne trzeba było dopuścić, żeby 12-bitowy przetwornik ADC i wbudowane kanały logiczne trafiły do oscyloskopów tańszych niż 10 000 zł. Rozdział 5 odpowiada na to pytanie dwuetapowo. Najpierw ustalamy punkt odniesienia – co realnie było dostępne około 2020 r. i za jakie pieniądze (5.1) – oraz co technicznie umożliwiło zejście tej technologii w dół (5.2). Następnie rozdzielamy marketing od rzeczywistości: ile naprawdę daje 12 bit rozdzielczości pionowej (5.3) i co wnosi architektura MSO (5.4). Kompromisy wymuszone ceną oraz przegląd konkretnych rodzin 2020...2026 zamkną rozdział w kolejnych sekcjach.

5.1. Punkt wyjścia: co realnie oferował rynek około 2020 r.

Wbrew powszechnemu wrażeniu 12-bitowy oscyloskop nie jest zdobyczą ostatnich dwóch-trzech lat. Natywne 12-bitowe DSO zadebiutowało 22 października 2012 r., kiedy Teledyne LeCroy przedstawił rodziny HDO4000 i HDO6000 – reklamowane jako pierwsza na rynku rodzina 12-bitowych oscyloskopów High Definition, opartych na technologii HD4096. Względem klasycznego 8-bitowego DSO dawały one szesnastokrotnie więcej poziomów kwantyzacji (4096 zamiast 256), przy paśmie od 200 MHz do 1 GHz. Problem leżał w cenie: modele HDO4000 kosztowały w katalogu od ok. 10 100 do 19 600 USD. To był sprzęt laboratoryjny, nie sprzęt na biurko pojedynczego konstruktora.

Podobnie było z MSO (*Mixed Signal Oscilloscope*), czyli oscyloskopem z wbudowanymi kanałami logicznymi. Sama idea sięga lat 90. (patrz oś czasu w rozdziale 1), a wersje mixed-signal HDO4000-MS oferowały 16 kanałów cyfrowych już w 2012 r. – ale zintegrowanie natywnej 12-bitowej akwizycji i toru logicznego w jednym, takim przyrządzie pozostawało rzadkością. Kto około 2020 r. chciał mieć naraz natywne 12 bit i kanały logiczne, płacił kilkadziesiąt tysięcy złotych, sięgając po klasę LeCroy HDO, Keysight S-series czy R&S RTO z trybem HD.

W segmencie budżetowym – poniżej 10 000 zł – dominował natomiast klasyczny 8-bitowy ADC. Pierwszym realnym krokiem w stronę wyższej rozdzielczości w tej klasie był R&S RTB2000 z 2017 r. z 10-bitowym przetwornikiem (opisany w rozdziale 1), ale **natywne 12 bit w segmencie budżetowym po prostu jeszcze nie istniało**. Teza tego rozdziału brzmi zatem prosto: między 2022 a 2024 r. ten próg runął – i to jest właściwa treść „zmiany 2020...2026”.

5.2. Co technicznie umożliwiło zejście 12 bit i MSO w dół

Skoro natywne 12 bit istniało od 2012 r., trzeba wyjaśnić, dlaczego dopiero po 2022 r. stało się tanie. Złożyły się na to trzy niezależne procesy. Pierwszy to dojrzałość i potanie przetworników ADC o architekturze pipelined (rozdział 2.4). To właśnie ta architektura – a nie flash – pozwala uzyskać 10...14 bitów przy próbkowaniu rzędu setek MSa/s do kilku GSa/s bez budowania 4096 komparatorów na jednym chipie. Kiedy 12-bitowe rdzenie pipelined

Skąd wzięto się natywne 12 bit: HD4096 (2012) a tryb hi-res

Warto trzymać w głowie rozróżnienie z rozdziału 2.4. Natywne 12 bit oznacza, że przetwornik ADC fizycznie ma 4096 poziomów i daje pełną rozdzielczość od razu, przy pełnym paśmie i pełnej szybkości próbkowania. Tryb hi-res to co innego – to oversampling z uśrednianiem, który „dorabia” dodatkowe bity z 8-bitowego ADC kosztem redukcji pasma. Rodzina LeCroy HDO4000/6000 (2012) była pierwszą, która dała 12 bit sprzętowo, cały czas. Dzisiejsze budżetowe konstrukcje (Siglent SDS800X HD, Rigol DH0800) idą tą samą drogą – natywne 12 bit z przetwornika – i to odróżnia je pokoleniowo od klasycznych 8-bitowych DSO z trybem hi-res.



Rigol DH0814 (Źródło: NDN.com.pl)

stały się tanim, katalogowym komponentem, znikła najdroższa bariera.

Drugi proces to spadek kosztu wydajnych układów SoC i FPGA pełniących rolę silnika akwizycji i renderingu. To, co w oscyloskopach klasy inżynierskiej realizował dedykowany ASIC (*Keysight MegaZoom*, *Tektronix DPX* – rozdział 2.5), w budżetowym oscyloskopie 2023 r. wykonuje programowalny układ za ułamek dawnej ceny. Trzeci proces to potaniecie pojemnościowych ekranów dotykowych i pamięci – dzięki czemu głęboka pamięć akwizycji i dotykowy interfejs przestały być wyróżnikiem najwyższej półki.

Efekt sumaryczny: producenci azjatyccy – Siglent, Rigol, UNI-T – mogli około 2023 r. złożyć w jednym przyrządzie natywny 12-bitowy ADC, 16 kanałów logicznych, dotykowy ekran i głęboką pamięć w cenie, która jeszcze w 2015 r. była nie do pomyślenia. Rohde & Schwarz doszła do tej samej fali inną drogą – przez własny, dedykowany ASIC akwizycji (200 Gb/s), ten sam, który realizuje w pełni cyfrowy układ wyzwalania opisany w rozdziale 1. Wprowadzony we wrześniu 2022 r. R&S MXO 4 to wariant premium tego samego przełomu: natywne 12 bit, architektura o rozdzielczości pionowej do 18 bit w trybie HD, 4,5 mln przebiegów na sekundę – ale ze startową ceną 7600 euro. Innymi słowy, ta sama technologia zeszła w dół dwiema ścieżkami naraz: azjatycką, napędzaną ceną, i R&S, napędzaną integracją i wydajnością akwizycji.

5.3. 12 bitów w segmencie budżetowym: realny zysk kontra marketing

Teorii nie wyprowadzamy tu drugi raz – wzór na stosunek sygnału do szumu kwantyzacji, $SNR = 6,02 \cdot N + 1,76$ dB, oraz pojęcie ENOB (*Effective Number Of Bits*) omówiliśmy w rozdziale 2.4 (patrz też rysunek 9). Przypomnijmy tylko wynik liczbowy: 8-bitowy ADC daje teoretycznie ok. 50 dB SNR, a 12-bitowy ok. 74 dB. Różnica to **24 dB – cztery pełne bity – w teorii**, a nie kilkanaście decybeli, jak bywa upraszczane.

Praktyka jest jednak skromniejsza, bo liczy się nie liczba w karcie katalogowej, lecz ENOB. Typowe wartości to 7,0...7,5 bit dla 8-bitowych i 9,5...10,5 bit dla 12-bitowych przetworników klasy inżynierskiej – czyli realny zysk rzędu 2,5...3 efektywnych bitów, ok. 15–18 dB. Dobrą ilustracją jest sam R&S MXO 4: mimo natywnego 12-bitowego ADC producent podaje dla niego ENOB na poziomie 10 bitów. Innymi słowy, około dwóch bitów „zjada” szum toru analogowego, jitter zegara i nieliniowości – i to zjada je nawet w konstrukcji z najwyższej półki.

Stąd pułapka segmentu budżetowego: w tanim oscyloskopie sam 12-bitowy ADC jest już niedrogi, ale **to front-end, nie przetwornik, decyduje o tym, ile z 12 bitów realnie zobaczymy**. Jeśli wzmacniacz wejściowy szumi, a zegar próbkowania ma duży jitter, różnica między 12 a 10 bitów potrafi się całkowicie zatrzeć. Nieprzypadkowo



Rigol DHO924 (Źródło: NDN.com.pl)

Siglent w opisie SDS800X HD akcentuje niski poziom szumów własnych front-endu (dziedzictwo konstrukcyjne Teledyne LeCroy) – bo właśnie ten parametr, a nie hasło „12 bit”, przesądza o jakości pomiaru małych sygnałów. Do tego wątku – i do tego, co budżetowy 12-bitowy MSO oddaje w zamian względem klasy high-end – wrócimy w sekcji 5.5, opartej na porównaniu radarowym (rysunek 15).

5.4. MSO: kanały logiczne w standardzie – jeden przyrząd zamiast dwóch

Druga oś zmiany rynkowej, równorzędna z 12 bit, to upowszechnienie architektury MSO. MSO to oscyloskop z wbudowanym, zintegrowanym analizatorem stanów logicznych – typowo 16 kanałów cyfrowych – dzielącym z torem analogowym wspólny układ wyzwalania i wspólną oś czasu. Dzięki temu na jednym ekranie widzimy jednocześnie przebieg analogowy i stany logiczne, a dekodowanie protokołów (I²C, SPI, UART, CAN, LIN) może korzystać z linii cyfrowych. Dla konstruktora oznacza to jedno: nie trzeba kupować osobnego analizatora stanów, a korelacja „zbrocze zegara analogowo kontra poziom logiczny” jest natychmiastowa.

Historia ceny tej funkcji jest równie wymowna jak w przypadku 12 bit. Szesnaście kanałów logicznych w mixed-signalowym LeCroy HDO4000-MS (2012) było elementem przyrządu za kilkanaście tysięcy dolarów. W 2023 r. Siglent

SDS800X HD daje 2 lub 4 kanały analogowe plus 16 kanałów cyfrowych w cenie poniżej ok. 600 USD – ta sama zdolność mixed-signal spadła zatem z klasy „ponad 10 000 USD” do „poniżej 600 USD” w ciągu jednej dekady.

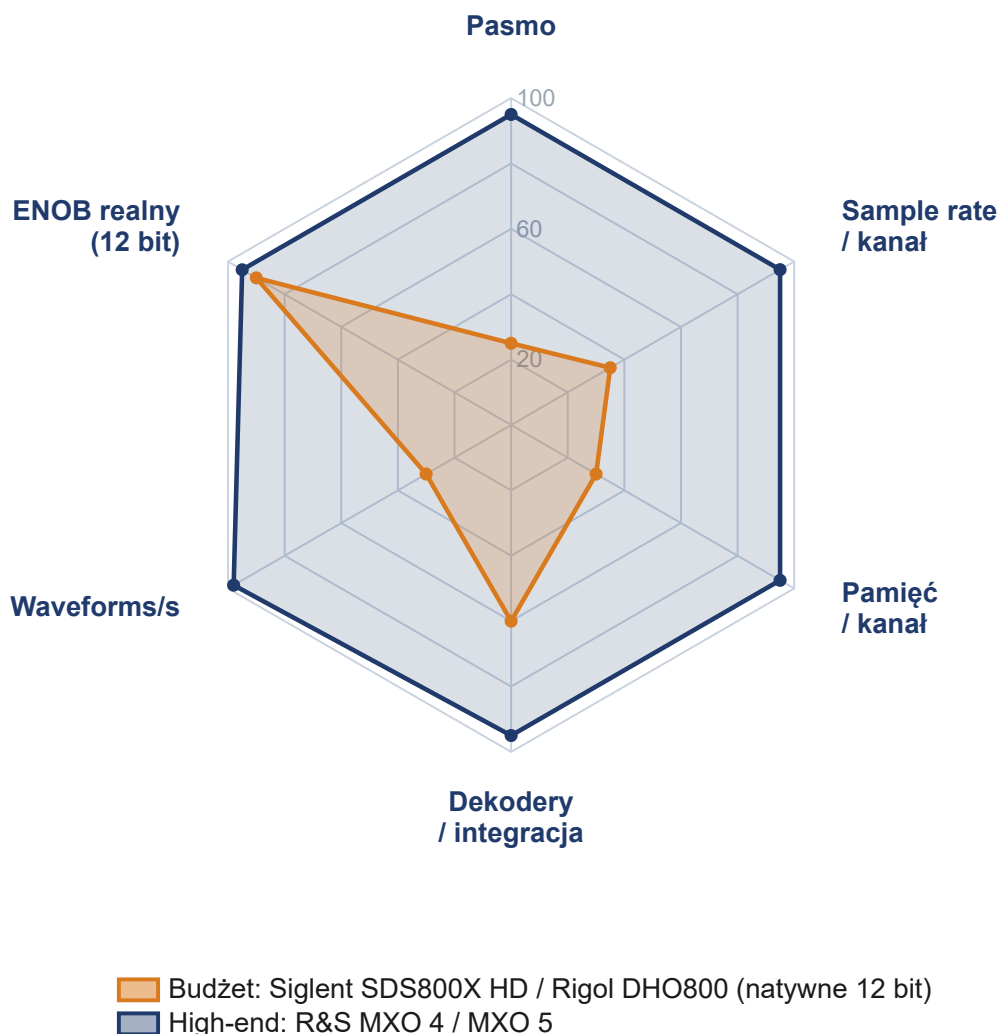
Rzetelność wymaga jednak dopowiedzenia, gdzie tanie MSO ustępują konstrukcjom droższymi. **Tor cyfrowy w budżetowym MSO ma zwykle niższą częstotliwość próbkowania i płytszą pamięć niż tor analogowy;** liczba kanałów logicznych to 16, nie 32; próg i histereza logiczna bywają wspólne dla całej grupy wejść, a nie ustawiane indywidualnie; wreszcie pełny zestaw dekodowników protokołów bywa płatną opcją, nie wyposażeniem standardowym. To realne ograniczenia, które trzeba zważyć przy zakupie – wracamy do nich w sekcji 5.5 (kompromisy) oraz w tabeli porównawczej 6.3.

5.5. Kompromisy wymuszone ceną – co budżetowy 12-bit MSO oddaje w zamian

Zejsście natywnego 12 bit i architektury MSO do segmentu budżetowego nie odbyło się za darmo. Żeby zobaczyć, za co realnie płacimy niższą ceną, warto zestawzić budżetowy 12-bitowy MSO (reprezentowany tu przez Siglent SDS800X HD i Rigol DHO800) z oscyloskopem klasy high-end (R&S MXO 4/5) w sześciu osiach naraz. Robi to porównanie radarowe na rysunku 15. Kluczowa obserwacja jest jedna: obie konstrukcje niemal pokrywają się tylko na jednej osi – rozdzielczości pionowej.

Budżetowy 12-bitowy MSO a oscyloskop high-end

sześć osi kompromisu (skala orientacyjna 0–100, na podstawie kart katalogowych)



Rysunek 15. Budżetowy 12-bitowy MSO (Siglent SDS800X HD/Rigol DHO800) a oscyloskop klasy high-end (R&S MXO 4/5) w sześciu osiach kompromisu; skala orientacyjna 0...100 w oparciu o karty katalogowe. Na osi rozdzielczości efektywnej (ENOB) budżetowy MSO niemal dorównuje high-endowi – to istota przetomu 12 bit. Przepaść pozostaje na paśmie, pamięci na kanał i szybkości akwizycji. Źródło: opracowanie graficzne EP na podstawie datasheetów Siglent SDS800X HD, Rigol DHO800/900 oraz R&S MXO 4/MXO 5

Oś rozdzielczości (ENOB) to jedyna, na której budżetowy MSO praktycznie dorównuje high-endowi. Natywny 12-bitowy ADC w połączeniu z porządnym front-endem daje realnie ok. 9,5...10 efektywnych bitów – podobnie jak 10-bitowy ENOB deklarowany dla R&S MXO 4. To właśnie sedno „rewolucji 12 bit”: rozdzielczość pionowa przestała być cechą odróżniającą klasy cenowe. Na pozostałych pięciu osiach przepaść jednak zostaje.

Pasma to największa różnica: budżetowe 12-bitowe MSO oferują typowo 70...250 MHz, podczas gdy high-end sięga 1,5...2 GHz. Kto mierzy szybkie zbrocza cyfrowe czy sygnały RF, w segmencie budżetowym po prostu nie ma pasma. Częstotliwość próbkowania na kanał jest drugą pułapką: konstrukcje budżetowe deklarują 1,25...2 GSa/s, ale ta wartość jest współdzielona między kanały przez interleaving (rozdział 3.4) i spada przy włączeniu wszystkich wejść. W klasie high-end każdy kanał ma własny,

dedykowany ADC pełnej szybkości (MXO: 5 GSa/s na kanał) i interleaving nie obowiązuje.

Pamięć na kanał rozjeżdża się podobnie: modele budżetowe mają 25...100 Mpts współdzielone między kanały, a R&S MXO 400...500 Mpts na każdym kanale jednocześnie – to decyduje o najdłuższym oknie obserwacji przy pełnym sample rate (rozdział 2.6). Szybkość odświeżania (waveforms/sec) w segmencie budżetowym mieści się w przedziale 120 tys.–1 mln na sekundę (milion w trybie Rigol UltraAcquire to wyjątek), podczas gdy MXO osiąga 4,5 mln – to realna różnica między „widzę rzadki glitch w ciągu minut” a „w ciągu sekund” (rozdział 2.5). Wreszcie integracja i dekodery: konstrukcje budżetowe dają podstawowe I²C/SPI/UART/CAN/LIN (część jako płatne pakiety), high-end dokłada rozbudowaną analizę widmową (45 000 FFT/s), wyzwalanie strefowe i długą listę protokołów.



Siglent SDS824X HD (Źródło: NDN.com.pl)

Wniosek dla konstruktora jest prosty: budżetowy 12-bitowy MSO to nie „tańszy high-end”, lecz **inny punkt na krzywej kompromisu** – wygrywa rozdzielczością pionową i ceną, przegrywa pasmem, pamięcią na kanał i szybkością akwizycji. Dobór właściwego oscyloskopu sprowadza się do dopasowania tych sześciu osi do własnej aplikacji, czym zajmiemy się w rozdziale 6.

5.6. Przegląd rodzin 2020...2026 – kto zepchnął 12 bit i MSO w dół

Przełom, o którym mowa, nie jest dziełem jednego producenta. Cztery rodziny konstrukcji – w różnych rolach – zdefiniowały rynek 2020...2026. Poniżej krótki przekrój każdej, z modelami reprezentatywnymi; szczegółowe zestawienie liczbowe znajdzie się w tabeli 6.3.

Siglent. Rodzina SDS800X HD (2023) to natywne 12 bit w cenie budżetowej: pasmo 70/100/200 MHz, 2 GSa/s, do 100 Mpts pamięci, 2 lub 4 kanały analogowe plus 16 cyfrowych, 120 000 wfm/s (500 000 w trybie sekwencyjnym), cyfrowy układ wyzwalania oraz dekodowanie magistral szeregowych. Wyżej w ofercie stoją SDS1000X HD, SDS2000X HD i SDS3000X HD – z większym pasmem i głębszą pamięcią. Atutem Siglenta jest niski poziom szumów front-endu (dziedzictwo konstrukcyjne Teledyne LeCroy), co – jak pokazaliśmy w 5.3 – realnie decyduje o tym, ile z 12 bit widać na ekranie. Pozycja: budżetowy lider jakości toru analogowego.

Rigol. Seria DHO800 (2023, platforma „Centaurus”) to najniższy próg wejścia do 12 bit – modele 70/100 MHz, 2 lub 4 kanały, 1,25 GSa/s, 25 Mpts, do miliona wfm/s w trybie

Tabela 6.3. Porównanie sześciu oscyloskopów reprezentatywnych dla 2026 r.

Model	Pasma	Kanały (A+C)	Sample rate	Pamięć	Rozdz.	MSO	Cena brutto*
Rigol DHO804	70 MHz	4	1,25 GSa/s	25 Mpts	12 bit	– (w DHO900)	~1 500...2 200 zł
Siglent SDS804X HD	70 MHz ¹	4 (+16)	2 GSa/s	50 Mpts	12 bit	opcja	~2 000...2 600 zł
UNI-T MSO1154HD	150 MHz	4 + 16	2,5 GSa/s	100 Mpts	12 bit	tak	~3 075 zł
Rigol DHO924	250 MHz	4 + 16	1,25 GSa/s	50 Mpts	12 bit	tak ²	~3 400 zł
R&S RTB2004	70...300 MHz	4	2,5 GSa/s	10 Mpts ³	10 bit	opcja	od ~9 200 zł netto
R&S MXO 44	200 MHz...1,5 GHz	4	5 GSa/s	400 Mpts/kan.	12 bit	opcja	od 7 600 €

¹ programowo odblokowywane do 200 MHz

² sonda logiczna PLA2216 osobno

³ do 160 Mpts z opcją historii/segmentacji

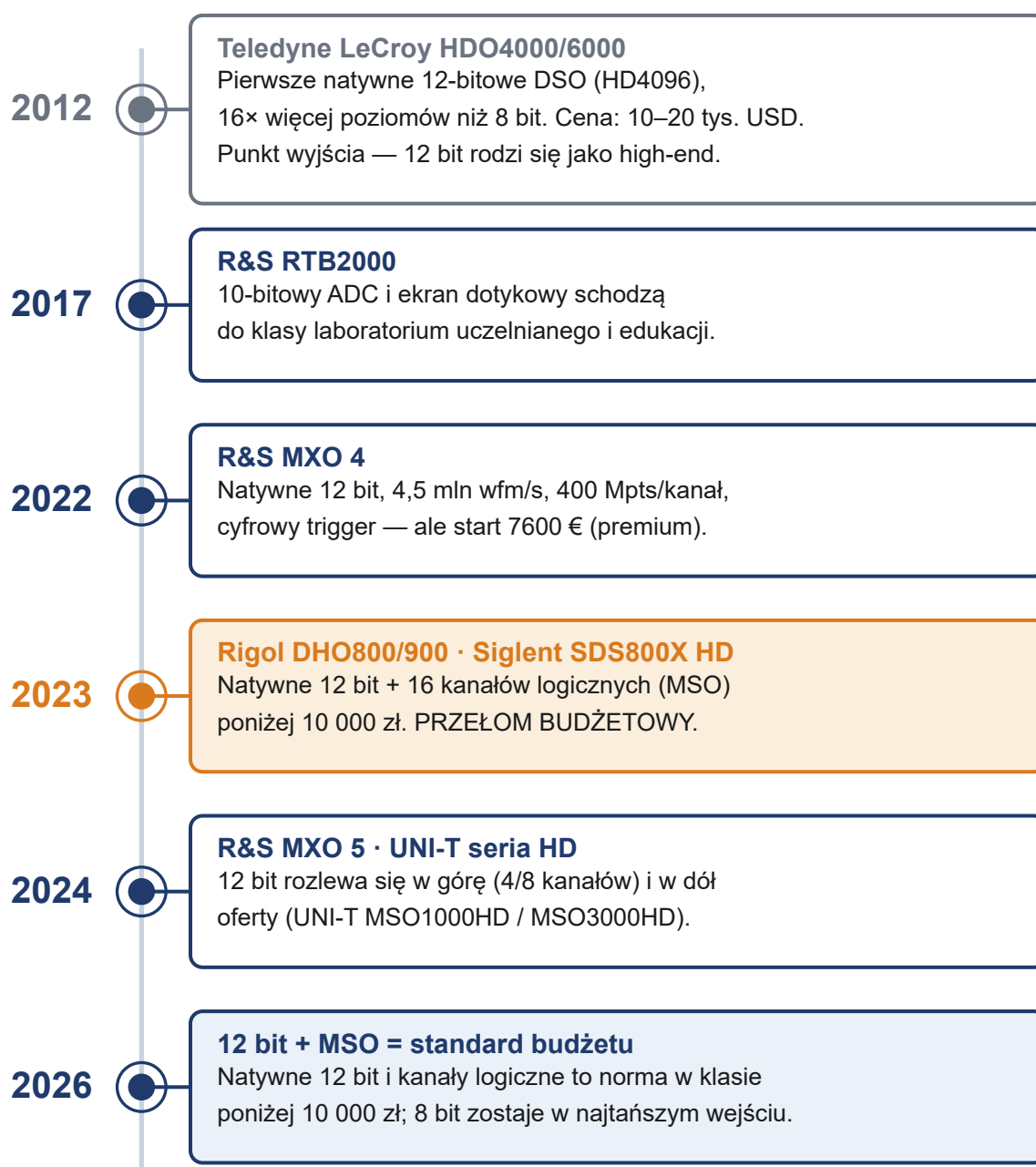
* Ceny orientacyjne, brutto (o ile nie zaznaczono „netto”), rynek PL, lipiec 2026;

źródła: siglent.com.pl, rigol.com.pl, uni-trend.pl, ndn.com.pl. R&S MXO 44 – cena katalogowa netto od 7600 €.

Sample rate podawany jest jako maksymalny, współdzielony między kanały (interleaving, rozdz. 3.4); przy pełnej liczbie kanałów spada

Jak 12 bit i MSO zeszły do klasy budżetowej

natywne 12-bitowe DSO: od high-endu (2012) do standardu budżetu (2026)



Rysunek 16. Oś czasu wejścia natywnej 12-bitowej rozdzielczości i architektury MSO do klasy budżetowej (2012...2026). Od high-endu (LeCroy HDO, 2012), przez R&S MXO 4 z klasy premium (2022), po przełom budżetowy 2023 (Rigol DHO800/900, Siglent SDS800X HD) i konsolidację 2024...2026. Źródło: opracowanie graficzne EP na podstawie materiałów Teledyne LeCroy (HD4096, 2012), Rohde & Schwarz (RTB2000 2017, MXO 4 2022, MXO 5 2024), Siglent (SDS800X HD, 2023), Rigol (DHO800/900, 2023) oraz UNI-T (seria HD)

UltraAcquire, zasilanie przez USB-C. Uzupełnia ją DHO900 (125/250 MHz, 4 kanały, 50 Mpts), która dodaje 16 kanałów cyfrowych, czyniąc z niej pełne MSO – z tym że sondę logiczną PLA2216 kupuje się osobno. To właśnie Rigol DHO800 „podpalił lont”, wymuszając na Siglencie szybką odpowiedź w postaci SDS800X HD. Pozycja: najtańsza droga do natywnego 12 bit i nowoczesnego interfejsu dotykowego.

UNI-T. Natywne 12 bit UNI-T oferuje w linii oznaczonej „HD” (MSO1000HD, MSO3000HD) – uwaga, starsze serie bez „HD” są jeszcze 8-bitowe. Reprezentatywny

MSO3024HD to 12 bit, 200 MHz, 4 kanały analogowe plus 16 cyfrowych, z wbudowaną analizą widmową, woltomierzem, częstotściomierzem i dekodernami – w duchu „dziewięć przyrządów w jednym”. Atutem jest agresywna integracja i cena; słabszą stroną – że część dekodernów i funkcji analitycznych to pakiety opcjonalne. Pozycja: maksimum funkcji w budżetowej cenie, kosztem dokupowanych opcji.

Rohde & Schwarz. R&S wchodzi w to zestawienie pełnym przekrojem, choć nie w roli najtańszego. Wejściem marki blisko segmentu budżetowego są RTB2000 (2017, 10-bitowy ADC, ekran dotykowy, 10 Mpts/kanal, 160 Mpts



Rigol DH0804 (Źródło: NDN.com.pl)

w trybie segmentowym) i RTM3000 (10 bit, głębsza pamięć). Natywne 12 bit R&S realizuje dopiero w serii MXO: MXO 4 (wrzesień 2022, pasmo 200 MHz...1,5 GHz, 4,5 mln wfm/s, 400 Mpts na każdym z czterech kanałów, cyfrowy trigger, start 7600 €) oraz MXO 5 (2024, 4 lub 8 kanałów, 12 bit, 5 GSa/s, 500 Mpts/kanał – pierwszy 8-kanałowy oscyloskop tej klasy). **Rzetelne pozycjonowanie: w segmencie 12-bitowym R&S gra w klasie premium** – wygrywa pasmem, pamięcią na kanał, szybkością akwizycji i integracją, ale na „oscyloskopie za złotówkę” w segmencie budżetowym Siglent i Rigol dają więcej sprzętu za tę samą kwotę. To rozgraniczenie – gdzie R&S wygrywa, a gdzie realnie przegrywa – pokażemy liczbowo w tabeli 6.3.

5.7. Podsumowanie: co naprawdę zmieniło się w latach 2020...2026

Historia, którą prześledziliśmy w tym rozdziale, układa się w spójną całość. Natywne 12 bit narodziło się w 2012 r. jako technologia high-end (Teledyne LeCroy HDO, cena rzędu 10...20 tys. USD) i przez dekadę tam pozostało. Punktem zwrotnym w segmencie premium był R&S MXO 4 (wrzesień 2022) – natywne 12 bit z dedykowanym ASIC-iem akwizycji, ale wciąż w cenie od 7600 €. Właściwy przełom budżetowy przyniósł rok 2023: Rigol DH0800/900 i Siglent SDS800X HD sprowadziły natywne 12 bit i architekturę MSO poniżej 10 000 zł. Lata 2024...2026 to konsolidacja trendu – R&S MXO 5, linia HD UNI-T

– po której 12-bitowy MSO stał się po prostu standardem klasy budżetowej.

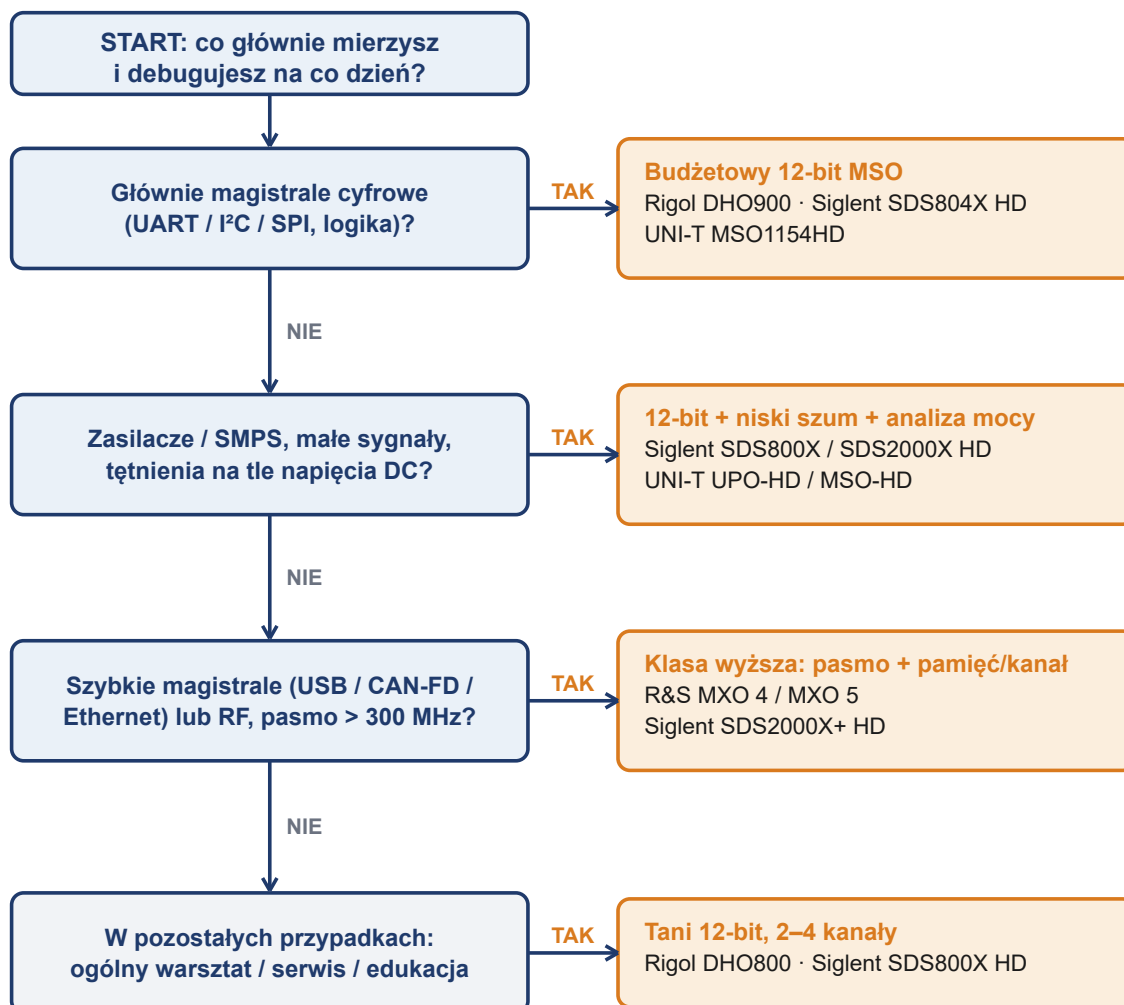
Dla konstruktora praktyczny sens tej zmiany sprowadza się do dwóch zdań. Po pierwsze, natywne 12 bit dają realnie ok. 2,5...3 dodatkowych efektywnych bitów rozdzielczości – ale tylko wtedy, gdy front-end (szum, jitter) na to pozwala; sam napis „12 bit” na obudowie niczego nie gwarantuje. Po drugie, MSO oznacza jeden przyrząd zamiast oscyloskopu i osobnego analizatora stanów. Ceną za tę dostępność są kompromisy z sekcji 5.5: ograniczone pasmo, płytsza pamięć na kanał, niższy sample rate na kanał i wolniejsza akwizycja. Oś czasu na rysunku 16 zbiera całą tę drogę w jednym widoku.

W rozdziale 6 przełożymy tę wiedzę na konkretny wybór: drzewo decyzyjne doboru oscyloskopu do typu pracy, szczegółową tabelę porównawczą modeli reprezentatywnych dla 2026 r. w kilku klasach cenowych, cztery typowe scenariusze pracy konstruktora oraz najczęstsze błędy zakupowe.

Rozdział 6. Praktyka konstruktora i porównanie modeli reprezentatywnych 2026

Rozdział 5 wyjaśnił, co i dlaczego zmieniło się na rynku. Ten rozdział przekłada tę wiedzę na jedną, praktyczną decyzję: który oscyloskop kupić. Zaczniemy od drzewa decyzyjnego (6.1), przejdziemy przez klasy cenowe 2026 r. (6.2) i tabelę porównawczą sześciu modeli reprezentatywnych (6.3),

Drzewo decyzyjne: jak dobrać oscyloskop do typu pracy



Rysunek 17. Drzewo decyzyjne doboru oscyloskopu do typu pracy. Ścieżka zaczyna się od pytania o dominujący rodzaj mierzonych sygnałów, a nie od nominalnego pasma; każda gałąź prowadzi do klasy sprzętu adekwatnej do zadania. Modele podane przykładowo, zgodnie z tabelą 6.3. Źródło: opracowanie EP

a następnie omówimy cztery typowe scenariusze pracy (6.4), niezbędne akcesoria (6.5) i najczęstsze błędy zakupowe (6.6).

6.1. Drzewo decyzyjne: dobór oscyloskopu do typu pracy

Najczęstszy błąd zaczyna się od złego pytania. Zamiast „ile pasma i za ile”, właściwe pierwsze pytanie brzmi: co głównie mierzę i debuguję na co dzień. Odpowiedź prowadzi do jednej z czterech typowych ścieżek, które zebrał na **rysunku 17**. Jeśli pracujesz głównie z magistralami cyfrowymi (UART, I²C, SPI), potrzebujesz przede wszystkim MSO i dekoderów, a nie gigahercowego pasma – wystarczy budżetowy 12-bitowy MSO. Jeśli mierzysz zasilacze, tętnienia i małe sygnały, kluczowe są rozdzielczość 12 bit, niski szum i analiza mocy. Dopiero szybkie magistrale (USB, CAN-FD, Ethernet) i sygnały RF wymagają realnie wyższego pasma i głębszej pamięci na kanał – czyli klasy droższej. W pozostałych zastosowaniach – ogólny warsztat, serwis, edukacja – najlepszym wyborem jest tani, natywnie 12-bitowy oscyloskop 2...4-kanałowy.

6.2. Klasy cenowe 2026 – co dostajesz za kolejne progi budżetu

Ceny na polskim rynku (lipiec 2026) układają się w kilka wyraźnych progów. Za około 1,5...3 tys. zł kupisz natywnie 12-bitowy oscyloskop 2...4-kanałowy o paśmie 70...200 MHz, z podstawowymi dekoderni i – zależnie od modelu – 16 kanałami logicznymi w standardzie lub opcji (Rigol DHO800, Siglent SDS800X HD, UNI-T MSO1154HD). To najlepszy stosunek możliwości do ceny w całej ofercie i punkt, w którym większość konstruktorów-amatorów i małych pracowni powinna zacząć.

Za około 3...5 tys. zł wchodzisz w pasmo 200...250 MHz, MSO w standardzie, głębszą pamięć i wyższą szybkość odświeżania (Rigol DHO924 ~3,4 tys. zł, UNI-T MSO1154HD ~3,1 tys. zł, wyższe warianty Siglent SDS800X/1000X HD). Za około 8...14 tys. zł masz dwie różne drogi: albo mocniejszy 12-bitowy budżet z dużym pasmem i pamięcią (Siglent SDS2000X HD), albo markowy, 10-bitowy przyrząd R&S RTB2004 (od ~9,2 tys. zł netto) z ekosystemem serwisowym i dojrzałym firmware, ale niższą rozdzielczością i płytszą pamięcią standardową niż tańszy



Siglent SDS814X HD (Źródło: NDN.com.pl)

Siglent. Powyżej 25...30 tys. zł zaczyna się klasa premium – R&S MXO 4 (od 7600 €) z natywnym 12-bitowym ADC, 4,5 mln przebiegów na sekundę, 400 Mpts na każdym kanale i pełną analizą widmową.

6.3. Tabela porównawcza modeli reprezentatywnych 2026

Zestawienie w tabeli 6.3 obejmuje sześć modeli rozłożonych na osi cenowej od ok. 1,5 tys. do ok. 40 tys. zł. Celowo zawiera R&S w dwóch punktach – 10-bitowy RTB2004 blisko środka i 12-bitowy MXO na szczycie – żeby pokazać, gdzie marka premium realnie wygrywa, a gdzie budżet daje więcej sprzętu za tę samą kwotę.

Najbardziej pouczające jest sąsiedztwo wierszy R&S i budżetowych. RTB2004 to solidny, markowy przyrząd z najlepszym w tabeli firmware i ergonomią, ale **jest 10-bitowy i za ok. 9,2 tys. zł netto oferuje mniejszą pamięć standardową niż 12-bitowy Siglent czy UNI-T za 2...3 tys. zł**. Z kolei MXO 44 wygrywa w tabeli niemal wszystkim – pasmem do 1,5 GHz, 5 GSa/s na kanał, 400 Mpts na kanał i 4,5 mln wfm/s – ale za cenę rzędu kilkudziesięciu tysięcy złotych. To jest właśnie sedno wyboru: nie „która marka najlepsza”, lecz „na której osi kompromisu (rysunek 15) leży moja aplikacja”.

6.4. Cztery scenariusze pracy konstruktora

(a) **Cyfrowy bring-up (UART/I²C/SPI)**. Tu liczy się MSO i dekodery, nie pasmo. Oscyloskop 70...200 MHz z 16 kanałami logicznymi i dekodowaniem magistral (Siglent SDS804X HD, Rigol DHO900, UNI-T MSO1154HD) w zupełności

wystarczy. Kluczowa jest korelacja czasowa zbrocza analogowego z liniami logicznymi – po to jest MSO (rozdział 5.4).

(b) **Pomiary zasilaczy i SMPS**. Rozdzielczość 12 bit ujawnia tętnienia na tle dużego napięcia stałego, których 8-bitowy przyrząd nie pokaże. Potrzebny jest niski szum front-endu, sondy różnicowe i prądowe oraz pakiet analizy mocy (Siglent SDS800X/2000X HD, UNI-T UPO/MSO-HD). To scenariusz, w którym natywne 12 bit daje najbardziej namacalny zysk.

(c) **Debugging szybkich magistral (USB/CAN-FD/Ethernet)**. Dopiero tutaj pasmo i pamięć na kanał stają się wąskim gardłem budżetu. Szybkie zbrocza i długie sekwencje pakietów wymagają pasma rzędu setek MHz i głębokiej pamięci utrzymywanej na wszystkich kanałach – czyli klasy R&S MXO lub Siglent SDS2000X+ HD.

(d) **Pomiary RF i małych sygnałów**. Liczą się FFT/analiza widma, uśrednianie i niski szum. Budżetowe 12-bitowe oscyloskopy mają sprzętowe FFT (do 1...2 Mpts), które wystarczą do diagnostyki EMI i harmonicznym; na szczycie R&S MXO wykonuje 45 000 FFT/s, co jest już inną ligą analizy widmowej.

6.5. Akcesoria niezbędne do każdej klasy

Cena oscyloskopu to nie cała cena pomiaru. Sondy pasywne są zwykle w zestawie, ale poważna praca wymaga dokupienia sond specjalistycznych: **różnicowych (pomiary SMPS i sieci), prądowych (pobór i przebiegi prądu) oraz aktywnych (szybkie zbrocza)**. W oscyloskopach MSO osobnym kosztem bywa sonda logiczna – np. Rigol PLA2216 czy UNI-T UT-M26 (ok. 1,9 tys. zł) – bez której 16 kanałów cyfrowych pozostaje

Głos praktyków

Poza kartami katalogowymi warto posłuchać, co powtarza się w środowisku – na forach branżowych i w publicznych recenzjach (m.in. EEVblog). Poniżej syntetyczny, sprawdzalny konsensus, a nie pojedyncze opinie.

- **„Pasma z okładki” to najczęstszy błąd.** Praktycy zgodnie wskazują, że o użyteczności decydują sample rate i pamięć na kanał, nie sama liczba MHz.
- **Przesunięcie punktu odniesienia.** Rigol DH0800 (pełna recenzja wideo EEVblog #1566, 2023) i Siglent SDS800X HD są uznawane za nowy standard „ile oscyloskopu za złotówkę” w budżecie.
- **Front-end ma znaczenie.** W testach porównawczych Siglent SDS800X HD bywa wskazywany jako mający niższy szum własny (rzędu 1...2 dB) niż budżetowe Rigole – co potwierdza tezę z 5.3, że o realnym zysku 12 bit decyduje tor analogowy, nie sam ADC.
- **Elastyczność i firmware.** Budżetowe Rigole i Siglenty mają programowo odblokowywane pasmo (np. 70→200 MHz); społeczność szeroko dokumentuje też modyfikacje firmware – na własne ryzyko i z utratą gwarancji.
- **Ergonomia bywa ważniejsza niż 100 MHz.** Dojrzałość oprogramowania i wygoda obsługi wracają w dyskusjach częściej niż nominalne parametry.

tylko zapisem w karcie katalogowej. Nie wolno też lekceważyć dobrej masy: jak pokazaliśmy w rozdziale 4, długi przewód masy potrafi zrujnować pomiar szybkiego zbocza niezależnie od klasy przyrządu.

6.6. Najczęstsze błędy zakupowe

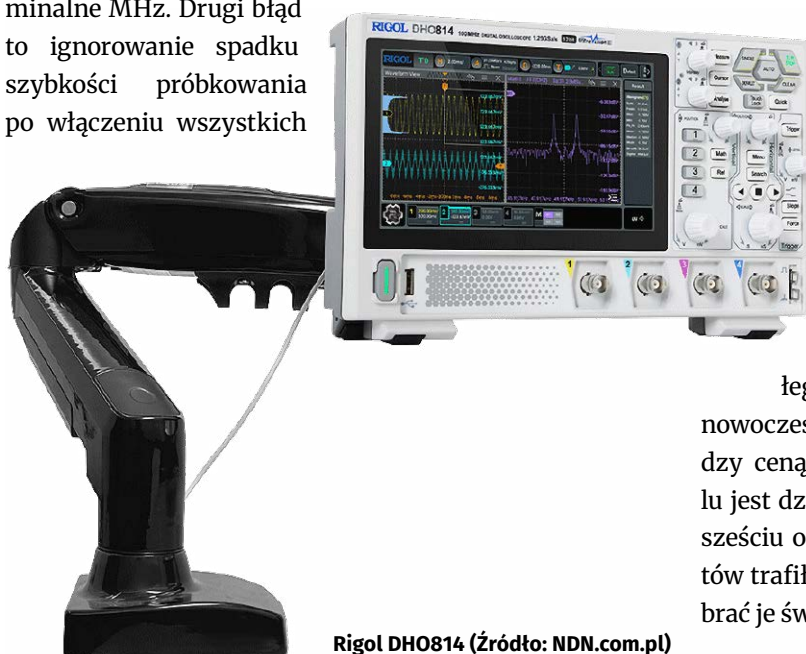
Pierwszy i najdroższy błąd to **kupowanie „pasma z okładki” zamiast sprzętu pod aplikację**. Dla większości zadań cyfrowych i zasilaczowych o użyteczności decydują sample rate na kanał i pamięć na kanał, a nie nominalne MHz. Drugi błąd to ignorowanie spadku szybkości próbkowania po włączeniu wszystkich

kanałów (interleaving, rozdział 3.4). Trzeci – pominięcie kosztu sond (logicznej, różnicowej, prądowej), który potrafi dodać do rachunku kilkadziesiąt procent. Czwarty – myślenie natywnego 12 bit z trybem hi-res (rozdział 5.1): tryb hi-res „dorabia” bity kosztem pasma, natywny 12-bitowy ADC daje je od razu. Piąty – przeplacanie za pasmo, którego nigdy nie wykorzystasz, przy jednoczesnym niedoszacowaniu pamięci, która realnie ogranicza długość obserwacji.

6.7. Podsumowanie całego cyklu

Przeszliśmy długą drogę – dosłownie. W części 1 zaczęliśmy od lampy Brauna i pierwszych oscyloskopów, w części 2 zeszliśmy na poziom próbkowania, wyzwalań, akwizycji i sond, a w tej, trzeciej, pokazaliśmy, jak natywne 12 bit i architektura MSO zeszły z laboratoryjnego high-endu do klasy budżetowej – i co to realnie oznacza przy zakupie. Jeśli z całego cyklu miałaby zostać jedna myśl, brzmi ona tak: nowoczesny oscyloskop przestał być kompromisem między ceną a rozdzielczością, a wybór właściwego modelu jest dziś nie kwestią marki, lecz trafnego dopasowania sześciu osi kompromisu do własnej pracy. Dwanaście bitów trafiło na biurko konstruktora – teraz wystarczy wybrać je świadomie.

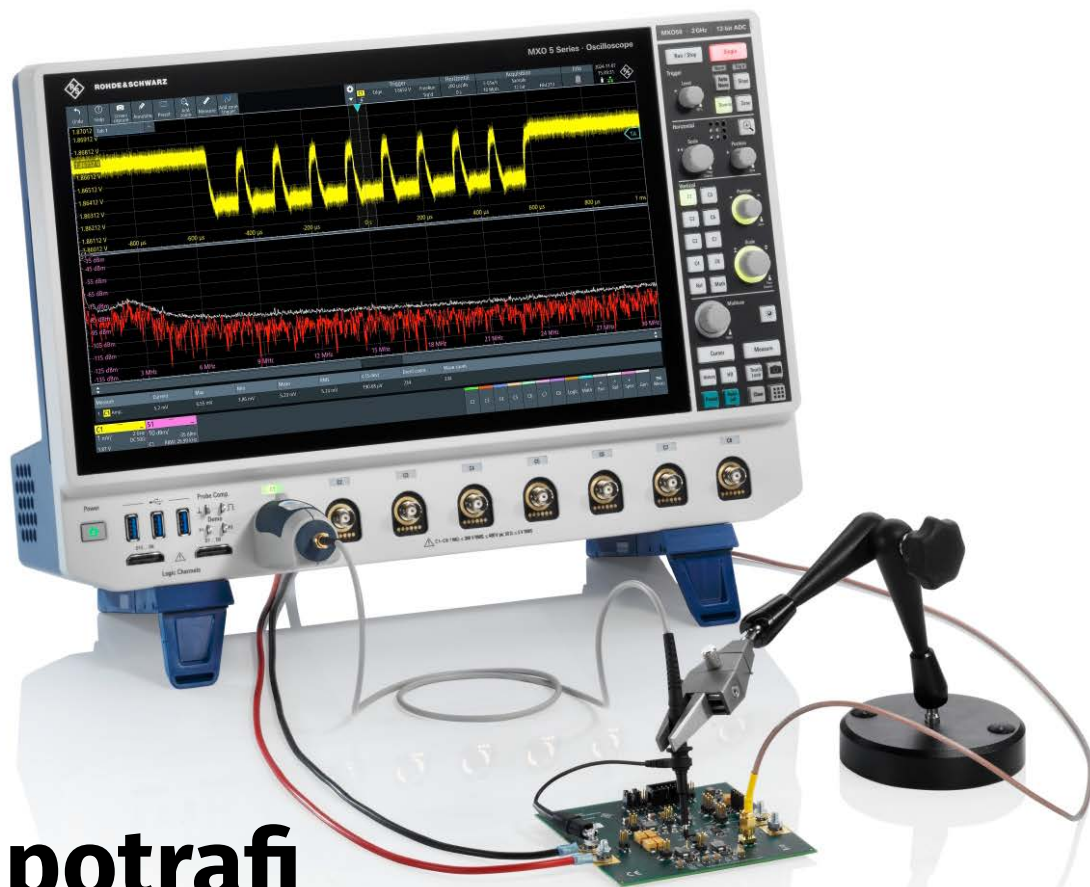
Wiesław Marciniak



Rigol DH0814 (Źródło: NDN.com.pl)

Bibliografia

- [1] Teledyne LeCroy, „HDO4000/HDO6000 High Definition Oscilloscopes” – technologia HD4096, komunikat i noty techniczne, 2012.
- [2] Rohde & Schwarz, „R&S®MXO 4 Series Oscilloscope” – karta katalogowa i komunikat prasowy, 2022.
- [3] Rohde & Schwarz, „R&S®MXO 5 Series Oscilloscope” – karta katalogowa, 2024.
- [4] Rohde & Schwarz, „R&S®RTB2000 Digital Oscilloscope” – karta katalogowa.
- [5] Siglent, „SDS800X HD Series Digital Oscilloscope” – karta katalogowa.
- [6] Siglent, „SDS2000X HD Series Digital Oscilloscope” – karta katalogowa.
- [7] Rigol, „DHO800/DHO900 Series (platforma Centaurus)” – karta katalogowa.
- [8] UNI-T, „MSO/UPO 1000HD i 3000HD Series” – karty katalogowe.
- [9] Teledyne LeCroy, „Understanding Oscilloscope Vertical Resolution and HD Technology” – nota techniczna.
- [10] Keysight, „Oscilloscopes Buying Guide” oraz „How to Select Your Next Oscilloscope (12 Tips)” – noty aplikacyjne.
- [11] Tektronix, „Oscilloscope Selection Guide” – przewodnik doboru.
- [12] Distrelec, „Przewodnik kupującego: wybór właściwego oscyloskopu”.
- [13] EEVblog #1566, „Rigol DH0800 12-bit Oscilloscope – Full Review”, 2023 (publiczna recenzja wideo; przywołana w ramce „Głos praktyków”).
- [14] Ceny detaliczne, rynek PL, lipiec 2026: siglent.com.pl, rigol.com.pl, uni-trend.pl, ndn.com.pl.
- [15] W. Marciniak, „Oscyloskopy 2026 – od lampy Brauna do 12 bitów na biurku konstruktora”, cz. 1–2, Elektronika Praktyczna 05–06/2026 (rozdziały 1–4).



Co potrafi nowoczesny oscyloskop

Sześć funkcji, które w praktyce zmieniają pracę inżyniera

Współczesny oscyloskop Rohde & Schwarz. Opisane niżej funkcje są dostępne w ofercie firmy – od serii RTB po rodziny MXO 4 i MXO 5. Materiały: Rohde & Schwarz.

W tym numerze zamykamy trzyczęściowe repetytorium o oscyloskopach. Pokazaliśmy w nim, że współczesny przyrząd to nie dawny oscylograf z wiązką, lecz komputer pomiarowy, a o jego przydatności decyduje nie samo pasmo, lecz zestaw funkcji wykonujących za inżyniera pracę dawniej żmudną albo wręcz niemożliwą.

Rohde & Schwarz udostępnił nam sześć przykładów takich funkcji, zilustrowanych zrzutami z własnych oscyloskopów. Poniżej krótko wyjaśniamy, czemu każda z nich służy i kiedy realnie skraca drogę do wyniku.

Głęboka pamięć akwizycji

Głęboka pamięć pozwala zarejestrować długie okno czasu bez rezygnacji z rozdzielczości. Bez niej trzeba wybierać: albo długi zapis przy niskiej częstotliwości

próbki, albo pełna szybkość na krótkim odcinku. Pamięć liczona w setkach milionów punktów znosi ten kompromis – jednym uchwyceniem obejmuje się na przykład całą sekwencję rozruchu układu przy pełnej szybkości próbkowania, a interesujące zdarzenie wyszukuje się dopiero potem, powiększając zapis. W rodzinie MXO pamięć sięga 400 milionów punktów na kanał, i to na wszystkich kanałach jednocześnie.

Tryb HD – czysty, odszumiony sygnał

Tryb HD (High Definition) podnosi efektywną rozdzielczość pionową przez nadpróbkowanie i filtrację, redukując szum toru wejściowego. Przebieg, który w trybie zwykłym tonie w szumie, w trybie HD staje się gładki i mierzalny – łatwiej dostrzec małe tętnienie na tle dużego napięcia



Rysunek 1. Głęboka pamięć: jedno uchwycenie obejmuje długie okno czasu i pozwala po fakcie wejść w dowolny szczegół sygnału bez ponownego pomiaru

stałego czy subtelne zniekształcenie. W oscyloskopach R&S rozdzielczość pionowa sięga w tym trybie 18 bitów.

Zone trigger – wyzwalanie, które się rysuje

Zone trigger zamienia żmudny dobór warunku wyzwalania na proste narysowanie obszaru na ekranie.

Zaznacza się strefę, a oscyloskop wyzwała tylko wtedy, gdy przebieg do niej wejdzie albo – przeciwnie – jej unika. Pozwala to wyłapać rzadkie zdarzenie lub odróżnić dwie zbliżone sytuacje bez definiowania złożonych warunków liczbowych. Funkcja działa zarówno w dziedzinie czasu, jak i w widmie, więc można wyzwać także na konkretnym zdarzeniu częstotliwościowym.

REKLAMA

Zestawy promocyjne pełne korzyści

Zaoszczędź do 56% na najpopularniejszych konfiguracjach aparatury pomiarowej



ROHDE & SCHWARZ
Make ideas real

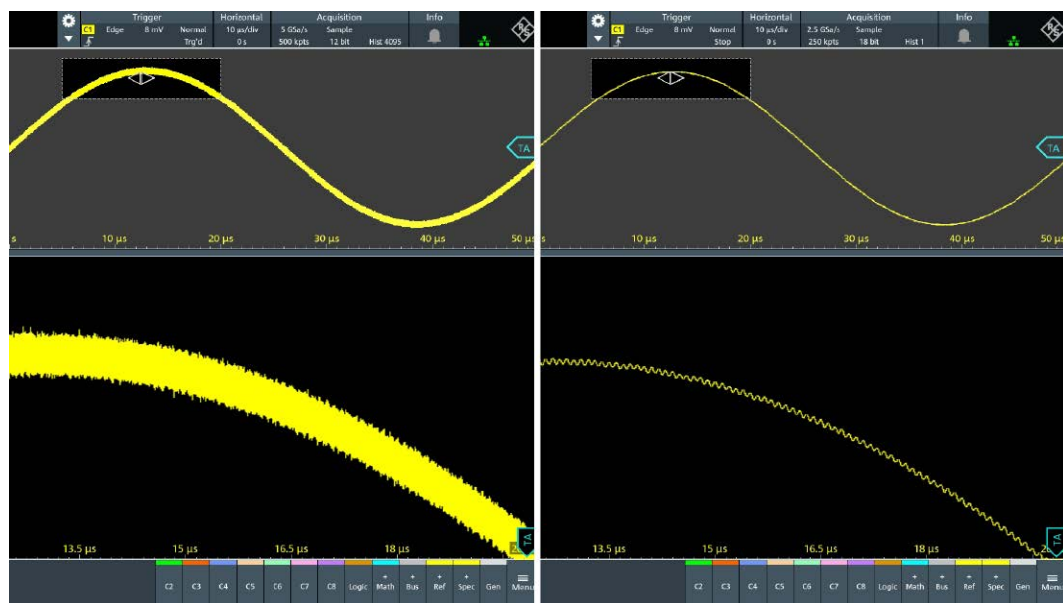
Rohde & Schwarz Polska Sp. z o.o.
00-024 Warszawa, Al. Jerozolimskie 44
tel. +48 662 106 936, www.rohde-schwarz.com



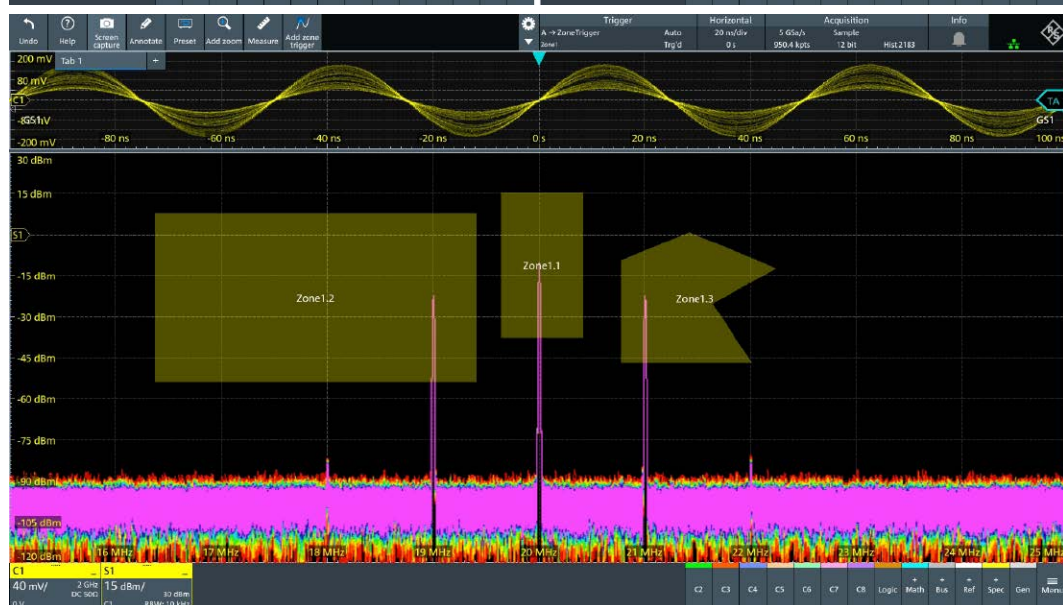
Analiza protokołów komunikacyjnych

Coraz więcej sygnałów przenosi się w domenę cyfrową, a dane biegną w ramach, które muszą trafiać na czas i bez

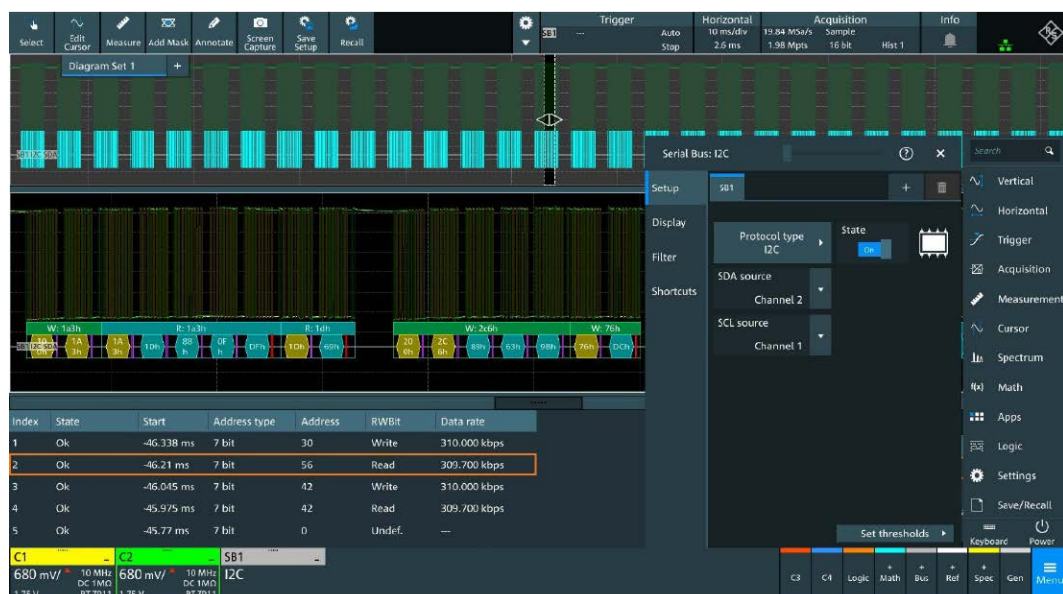
błędu. Wbudowane dekodery rozkładają magistrale – I²C, SPI, UART, CAN i inne – na czytelne ramki, zsynchronizowane z przebiegiem elektrycznym i zebrane w tabeli.



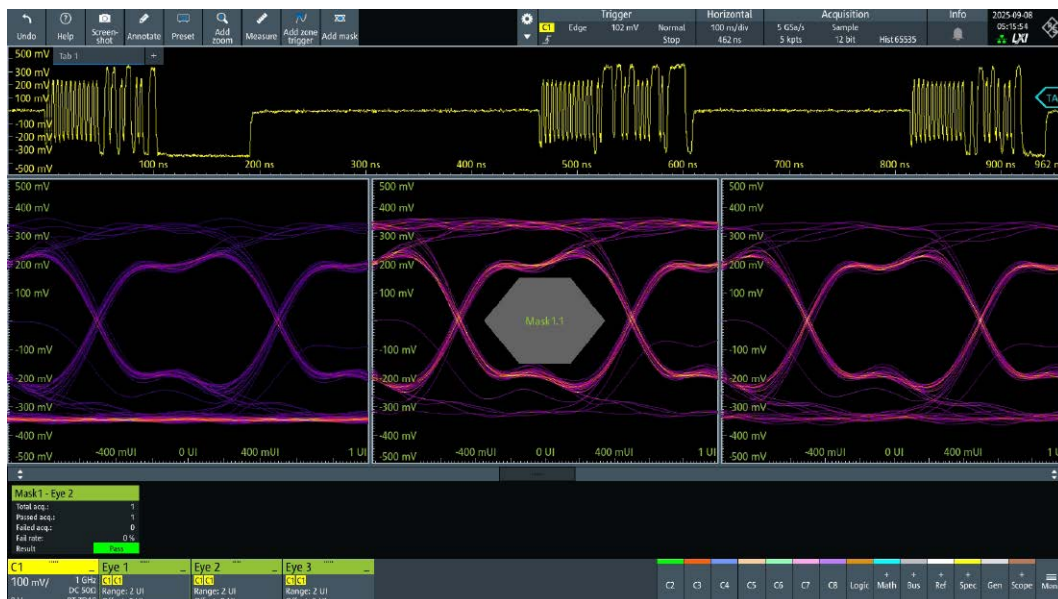
Rysunek 2. Tryb HD: ten sam sygnał bez i z odszumianiem. Wyższa rozdzielczość pionowa ujawnia szczegóły ukryte w szumie



Rysunek 3. Zone trigger: wyzwalanie przez wskazanie obszaru na ekranie, tu w połączeniu z analizą widmową



Rysunek 4. Dekodowanie protokołu: ramki magistrali zestawione w tabeli i powiązane z przebiegiem elektrycznym



Rysunek 5. Diagram oka sygnału USB 2.0: otwarcie „oka” mówi o marginesie na szum i jitter

Dzięki temu jednocześnie widać zawartość logiczną transmisji i jej stronę analogową; łatwiej wychwycić błąd ramkowania, zły poziom logiczny czy zakłócenie na linii.

Signal integrity i diagram oka

Przy bardzo szybkich sygnałach cyfrowych o poprawności transmisji decyduje integralność sygnału (signal integrity). Nałożenie na siebie wielu okresów bitowych daje diagram oka (eye diagram): im „oko” szerzej otwarte, tym większy margines na szum i jitter, a im bardziej przymknięte – tym bliżej błędów. To standardowe narzędzie oceny łączy szybkich, na przykład USB, gdzie jednym spojrzeniem widać, czy sygnał mieści się w wymaganej masce.

Charakterystyka Bodego

Z wbudowanym generatorem oscyloskop może wykreślić charakterystykę Bodego – wzmacnienie i fazę w funkcji częstotliwości. To narzędzie do analizy filtrów, elementów pasywnych, pętli sprzężenia wzmacniaczy operacyjnych i stabilności zasilaczy. Wątek łączy się wprost z artykułem o stabilności wzmacniaczy operacyjnych w tym numerze: charakterystyka Bodego pokazuje zapas fazy, zanim układ zacznie oscylovować, i pozwala dobrać kompensację na podstawie pomiaru, a nie założeń.



Rysunek 6. Charakterystyka Bodego: wzmacnienie i faza w funkcji częstotliwości – podstawa oceny stabilności i doboru kompensacji

O produktach

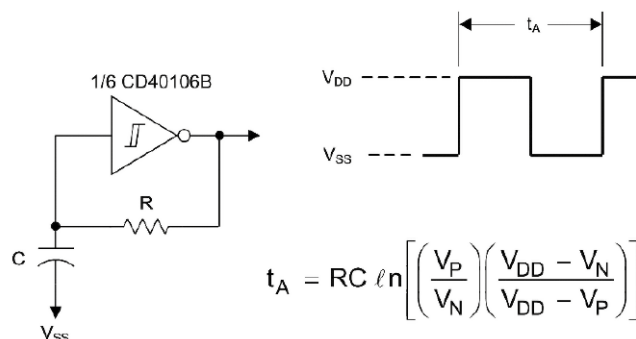
Wszystkie opisane funkcje są dostępne we współczesnych oscyloskopach Rohde & Schwarz – od kompaktowej serii RTB po rodziny MXO 4 i MXO 5 z natywną 12-bitową akwizycją, szybkim cyfrowym torem wyzwalań i rozbudowaną analizą widmową. Szczegółowe dane oraz materiały dla mediów udostępnia Rohde & Schwarz.

Astabilne układy czasowe

Generowanie sygnału prostokątnego to banalna rzecz dla dzisiejszych mikrokontrolerów i układów FPGA. Ale czy da się to zrobić prościej, bez angażowania tak drogiego podzespołu? Okazuje się, że tak, zaś możliwości jest znacznie więcej niż tylko topowy „555”.

Nie w każdym zastosowaniu potrzebna jest wysoka stabilność generowanej częstotliwości. Pulsująca lampa ostrzegawcza, przerywany sygnał dźwiękowy czy migająca dioda na obudowie urządzenia – te i inne aplikacje będą z powodzeniem działały prawidłowo przy nawet kilkustopniowym odchyleniu częstotliwości od swojej wartości nominalnej. Najprostszym rozwiązaniem jest wsadzenie w układ „procka” i nakazanie mu cyklicznej zmiany stanu logicznego na jednym z wyprowadzeń, ale czy nie można prościej?

Owszem, można – nawet na trzech elementach. Mowa o układzie z rysunku 1, który zawiera rezystor, kondensator i odwracającą bramkę logiczną z wejściem Schmitta. Wprawdzie w tej bramce znajduje się znacznie więcej podzespołów, ale dla nas nie ma to większego znaczenia – mamy oscylator jak ta lala. Stosując miniaturowe



Rysunek 1. Schemat multiwibratora astabilnego z bramką NOT oraz wzór na okres sygnału wyjściowego [1]

elementy oraz bramkę w obudowie SOT23-5, jak na przykład SN74LVC1G14, można uzyskać bardzo przyjemny oscylator RC. Zasada działania tego układu nie jest zbyt skomplikowana: rezystor ładuje kondensator prądem wpływającym z wyjścia bramki, przez co napięcie na tym kondensatorze rośnie. Kiedy urośnie do wartości odpowiadającej górnemu progowi przerzutu, wyjście bramki zmienia swój stan na niski i zaczyna się rozładowywanie kondensatora przez rezystor, aż do osiągnięcia dolnego progu przerzutu – wówczas cykl się zamyka.

W teorii można łatwo obliczyć częstotliwość oscylacji – odpowiedni wzór jest podany na rysunku 1. Praktyka jest przysparza więcej problemów, bowiem:

PARAMETER	TEST CONDITIONS	V _{CC}	-40°C to 85°C			-40°C to 125°C ⁽²⁾			UNIT
			MIN	TYP ⁽¹⁾	MAX	MIN	TYP	MAX	
V _{T+} Positive-going input threshold voltage		1.65 V	0.79		1.16	.79		1.16	V
		2.3 V	1.11		1.56	1.11		1.56	
		3 V	1.5		1.87	1.5		1.87	
		4.5 V	2.16		2.74	2.16		2.74	
		5.5 V	2.61		3.33	2.61		3.33	
V _{T-} Negative-going input threshold voltage	DBV, DCK, DRL, DRY, DSF, YZV and YZP packages	1.65 V	0.39		0.62	.39		.64	V
		2.3 V	0.58		0.87	.58		.89	
		3 V	0.84		1.14	.84		1.16	
		4.5 V	1.41		1.79	1.41		1.79	
		5.5 V	1.87		2.29	1.87		2.29	
V _{T-} Negative-going input threshold voltage	DPW package	1.65 V	0.44		0.67				V
		2.3 V	0.63		0.92				
		3 V	0.89		1.19				
		4.5 V	1.46		1.84				
		5.5 V	1.92		2.34				
ΔV _T Hysteresis (V _{T+} - V _{T-})		1.65 V	0.37		0.62	0.37		0.62	V
		2.3 V	0.48		0.77	0.48		0.77	
		3 V	0.56		0.87	0.56		0.87	
		4.5 V	0.71		1.04	0.71		1.04	
		5.5 V	0.71		1.11	0.71		1.11	

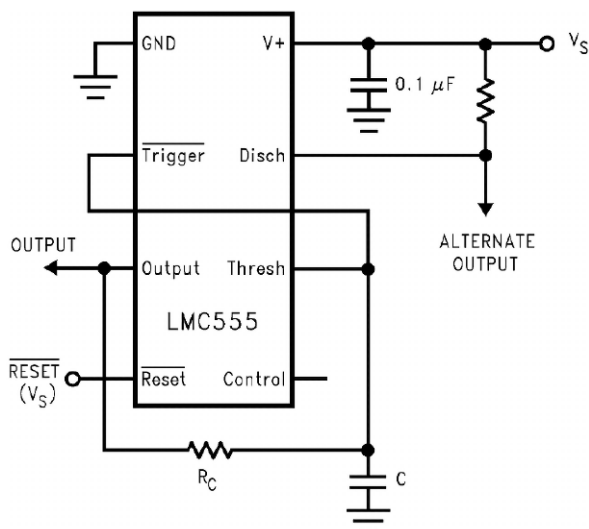
Rysunek 2. Progi przerzutu układu SN74LVC1G14 [2]

- pojemność kondensatorów, zwłaszcza miniaturowych ceramicznych X7R, silnie zależy od napięcia,
- napięcie wyjściowe przerzutnika nie osiąga wartości VDD i VSS z uwagi na prąd płynący przez rezystor,
- progi przerzutu przerzutnika są zdefiniowane z silnymi dopuszczalnymi rozrzutami.

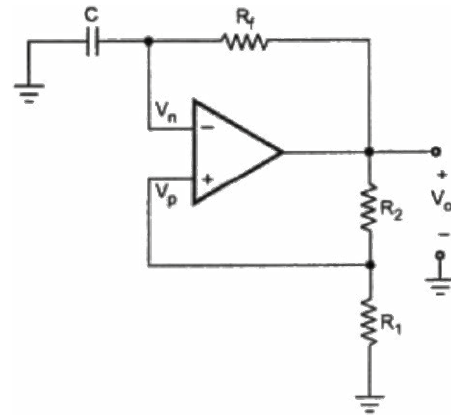
Ostatni punkt ma w praktyce największe znaczenie, zaś szczegółów dostarcza **rysunek 2**. Nie trzeba wyższej matematyki by przekonać się, że progi przerzutu tego układu mogą mieć nawet kilkadziesiąt procent rozrzutu. Z tego powodu takie oscylatory są przydatne w bardzo prostych aplikacjach, jak miganie diodami LED na panelu kontrolnym czy generowanie pisku przez buzzer piezoelektryczny, lecz na tym ich rola się zdecydowanie kończy. Jest jeszcze jedna duża zaleta: dobierając R1 o wysokiej rezystancji, na przykład setek kiloomów, można uzyskać naprawdę atrakcyjny pobór prądu – zapotrzebowanie układu SN74LVC1G14 nie przekracza 10 μ A.

Skoro progi przerzutu są problemem, to może warto sięgnąć po układ, który ma je całkiem dobrze zdefiniowane. Mam tutaj na myśli legendarnego „555”, który ma już grubo ponad 50 lat. I dobrze – to, co dobre, cały czas ma swoich odbiorców. Nie chcę tu jednak opisywać znanej do przesady aplikacji układu astabilnego z dwoma rezystorami i kondensatorem, ponieważ można zrobić to znacznie prościej – **rysunek 3**. Wystarczy użyć 555 jako inwertera z wejściem Schmitta, którego progi przerzutu są określone jako 1/3 i 2/3 napięcia zasilania. Rezystor dołączony do wyjścia DISCH nie jest konieczny do poprawnego działania tego układu, jedynie podciąga wyjście typu „otwarty kolektor” do dodatniej linii zasilania, co umożliwia uzyskanie dodatkowego wyjścia, będącego w przeciwfazie do wyjścia OUTPUT.

Producent w notce katalogowej podaje wersję LMC555, a to dlatego, że ma wejścia typu CMOS, które nie pobierają prądu, oraz wyjście typu CMOS, które z kolei osiąga potencjały możliwie zbliżone do napięcia zasilającego.



Rysunek 3. 555 w roli multiwibratora o wypełnieniu 50% [3]

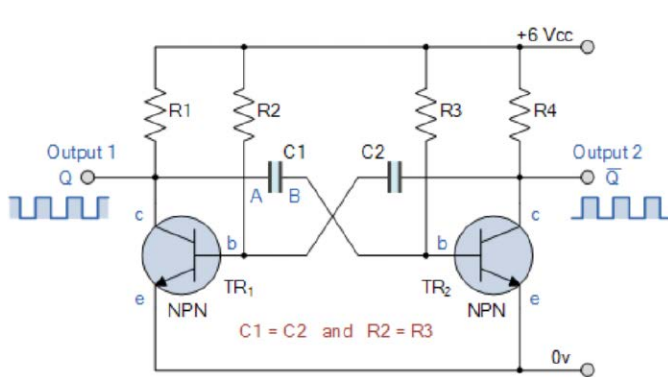


Rysunek 4. Najprostszy generator astabilny ze wzmacniaczem operacyjnym [4]

To pozwala uzyskać – teoretycznie – 50% wypełnienia sygnału wyjściowego. Stosując inną wersję, na tranzystorach bipolarnych, trzeba się liczyć z tym, że uzyskana częstotliwość może się różnić od wzorcowej: $f = 1/(1,4 \cdot RC \cdot C)$. Ten wzór jest niczym innym jak sprowadzeniem wzoru z rysunku 1 do postaci uproszczonej, przy znanych progach przerzutu. Moja doświadczenia z tym układem są bardzo pozytywne, osiągnięcie zakładanej częstotliwości oraz wypełnienia bliskiego 50% są naprawdę realne. W przypadku „zwykłych” bramek z wejściem Schmitta, wypełnienie sygnału nie będzie równe 50%, ponieważ progi przerzutu nie są ułożone w nich symetrycznie względem połowy napięcia zasilającego.

Można dalej powalczyć o progi przerzutu, bazując na układzie, który jest do tego specjalizowany. Mam na myśli wzmacniacz operacyjny w aplikacji przerzutnika Schmitta, jak na **rysunku 4**. Należy pamiętać, że w tym układzie wzmacniacz operacyjny jest zasilany napięciem symetrycznym, przy unipolarnym zasilaniu należy zastosować „sztuczną masę”. Stosując komparator, ale z wyjściem totem-pole, można uzyskać naprawdę strome zbocza sygnału wyjściowego, bowiem w typowych wzmacniaczach operacyjnych ograniczeniem jest ich slew-rate. Jeżeli tylko wyjście wzmacniacza tudzież komparatora przyjmuje symetryczne wartości (np. 13 V i –13 V), uzyskany przebieg będzie miał wypełnienie 50%. Można jednak ten układ całkiem ciekawie zmodyfikować, dodając ograniczniki lub przesuwniki napięcia wyjściowego, by mieć wpływ na wypełnienie.

We wszystkich wymienionych wcześniej multiwibratorach astabilnych za częstotliwość oscylacji odpowiadał tylko jeden rezystor i jeden kondensator. W takim zestawieniu nie można zabraknąć bardzo prostego, bardzo znanego i sprawiającego wiele kłopotów układu Ecclesa-Jordana. Już wjeżdża na **rysunek 5**, Proszę Państwa, oto on! Mój Promotor mawiał, że połowę tego układu zaprojektował Eccles, a drugą Jordan... Układ bardzo prosty w zrozumieniu, dopóki pozostaje na tablicy.

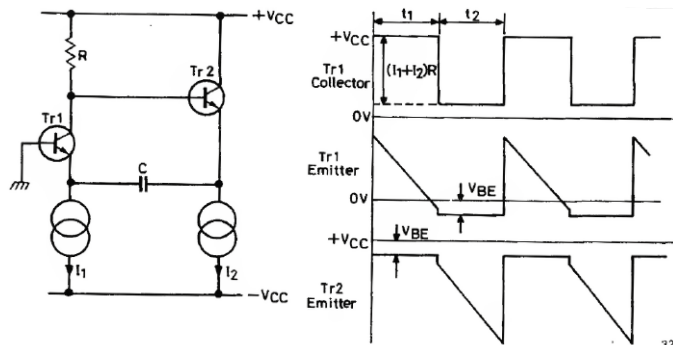


Rysunek 5. Układ multiwibratora Ecclesa-Jordana [5]

Jak to bywa z prostymi układami, tutaj wszystko ma wpływ na wszystko. Przebiegi na kolektorach wcale nie są prostokątne, bardziej „trapezoidalne”, a to przez fakt, że przez rezystor kolektora zatkanego tranzystora płynie prąd ładujący kondensator tranzystora naprzeciwko. Ale to nie wszystko. W chwili przerzutu, kiedy jeden tranzystor nagle się zatyka, a drugi gwałtownie nasycza, zbyt wysokie napięcie na okładkach naładowanego kondensatora jest w stanie przebić złącze baza-emiter zatkanego tranzystora w kierunku zaporowym, a to z kolei doprowadza do jego niekontrolowanego przewodzenia. Z tego powodu, na rysunku 5 znalazło się zasilanie napięciem 6 V, bowiem przy wyższym może dochodzić do tego właśnie efektu – pamiętam, ile to ja się natrudziłem za młodu, by ten układ uruchomić przy zasilaniu napięciem 24 V... To się nie uda. Rozwiązaniem mogą być tranzystory MOSFET, ale tu z kolei uwaga na wytrzymałość izolatora podbramkowego. Ten układ jest wrażliwy na... W sumie, to na wszystko, gdyż nawet rezystancja obciążenia ma wpływ na częstotliwość. Pomijając tak drobny fakt, że w ustalaniu częstotliwości biorą udział tak naprawdę cztery (!) i rezystory i dwa kondensatory. Często pokazuje się go dzieciom na podstawowych zajęciach z elektroniki, ale zjawiska w nim zachodzące to temat na spory wykład. Trzeba mu jednak oddać, że do pracy wystarczą tylko dwa elementy aktywne, co sto lat temu było niebagatelną zaletą. Można ten układ rozbudować o diody ograniczające napięci baza-emiter i nie tylko (szczegóły w [6]), ale – moim zdaniem – dzisiaj to już tylko sztuka dla sztuki.

Dla miłośników elektroniki analogowej, a sam do takich się zaliczam, przygotowałem na koniec mało znany smaczek: przerzutnik Bowesa. Jego schemat w najprostszym wydaniu znajduje się na **rysunku 6**. Jego wielką zaletą jest fakt, że żaden z tranzystorów nie wchodzi tutaj w stan nasycenia – jest zatkany lub w stanie aktywnym – więc można on osiągać bardzo wysokie częstotliwości przełączania, nieosiągalne dla wcześniejszych układów. Dwa tranzystory, rezystor, kondensator, dwa źródła prądowe... Jak to działa?

Załóżmy wstępnie, że Tr1 właśnie się zatkał, a Tr2 jest otwarty – jego bazę polaryzuje rezystor R. Tranzystor ten



Rysunek 6. Przerzutnik Bowesa z podstawowymi przebiegami [6]

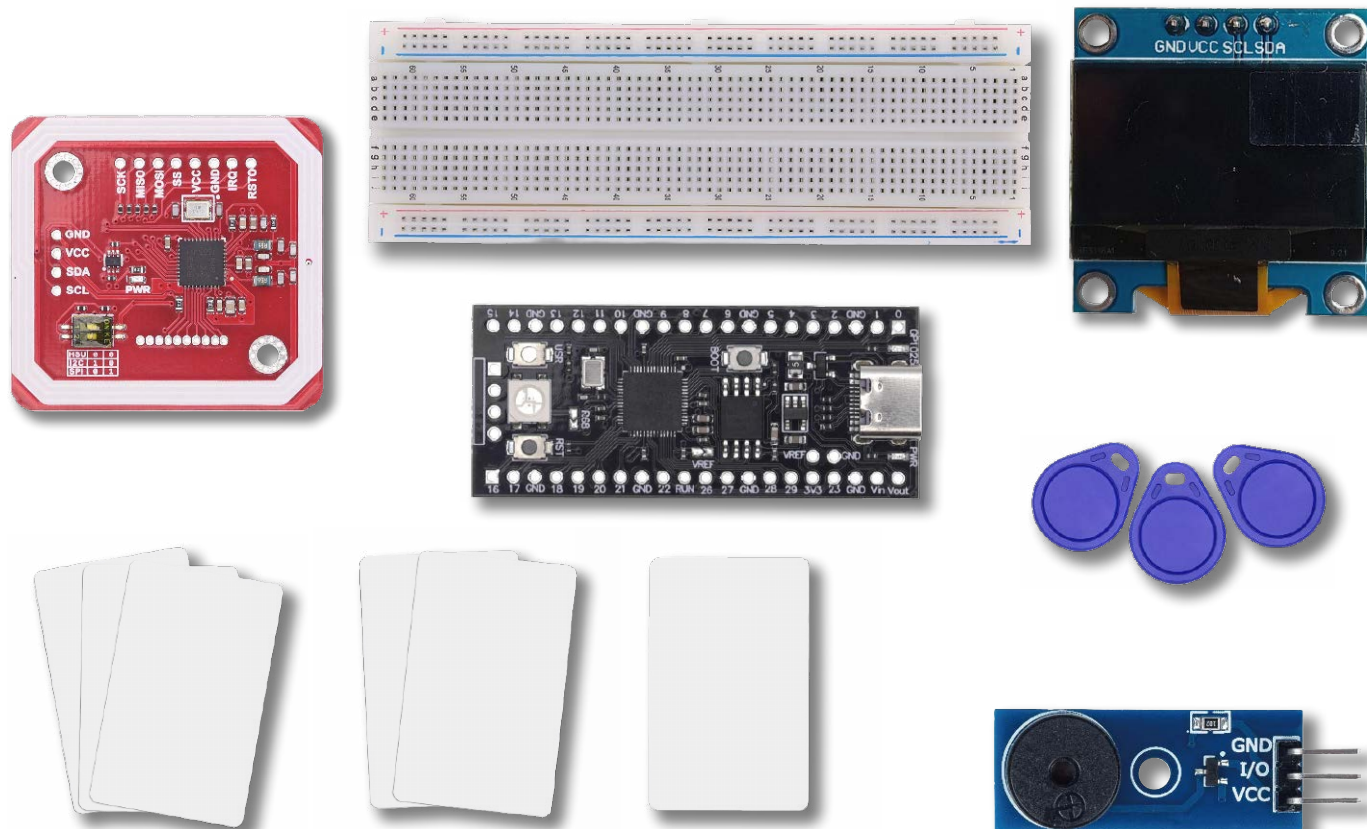
przewodzi swoim emiterem prąd źródła I2 oraz umożliwia kondensatorowi C liniowe ładowanie się poprzez źródło prądowe I1, które nie ma innej drogi przepływu prądu, jak tylko przez ten kondensator. W momencie, w którym kondensator C naładowuje się na tyle, że emiter Tr1 osiągnie potencjał $-0,7$ V, ulegnie on załączeniu. Wtedy potencjał kolektora Tr1 spada, bo występuje spadek napięcia na rezystorze R, zaś Tr2 zatyka się, spolaryzowany zaporowo. Przez emiter Tr1 płyną wtedy prądy obu źródeł prądowych, I1 i I2, podobnie jak wcześniej płynęły przez emiter Tr2. Teraz z kolei kondensator C przeładowuje się w drugą stronę, gdyż źródło I2 „wyciąga” z niego prąd w przeciwnym kierunku, niż poprzednio robiło to źródło I1. I znowu: kiedy C naładowuje się na tyle, by złącze baza-emiter Tr2 znów zaczęło przewodzić, sytuacja odwraca się i układ wraca do stanu początkowego.

Można dodać, że ten układ to jedno z najszybszych dostępnych nam rozwiązań, bowiem tranzystory w nim przechodzą ze stanu aktywnego wprost do zatkania, będąc nawet polaryzowane zaporowo, co przyspiesza ich wyłączenie. Warto zauważyć, że na wypełnienie sygnału mają wpływ wartości prądów I1 i I2, które to źródła prądowe mogą być – w najprostszym wydaniu – zwykłymi rezystorami. Szkoda, że ten układ dzisiaj jest już głównie ciekawostką, lecz lata temu pozwalał na uzyskiwanie częstotliwości generacji rzędu setek megaherców. Dlatego uznałem za stosowne wspomnieć o nim choćby w kilku akapitach.

Michał Kurzela, EP

Źródła:

- [1] <https://www.ti.com/lit/ds/symlink/cd40106b.pdf>
- [2] <https://www.ti.com/lit/ds/symlink/sn74lvc1g14.pdf>
- [3] <https://www.ti.com/lit/ds/symlink/lmc555.pdf>
- [4] https://entri.app/blog/wp-content/uploads/2022/07/HSSST_Electronics_Part3_Op_amp_based_multivibrators.pdf
- [5] <https://www.electronics-tutorials.ws/waveforms/astable.html>
- [6] <https://www.mikrocontroller.net/attachment/338989/astable.pdf>



Kurs RFID z PN532, rozdział 2

Norma ISO/IEC 14443-3: pętla antykolizyjna i identyfikacja kart

Pierwsza część
w EP06/2026 – dostępna na
www.UlubionyKiosk.pl

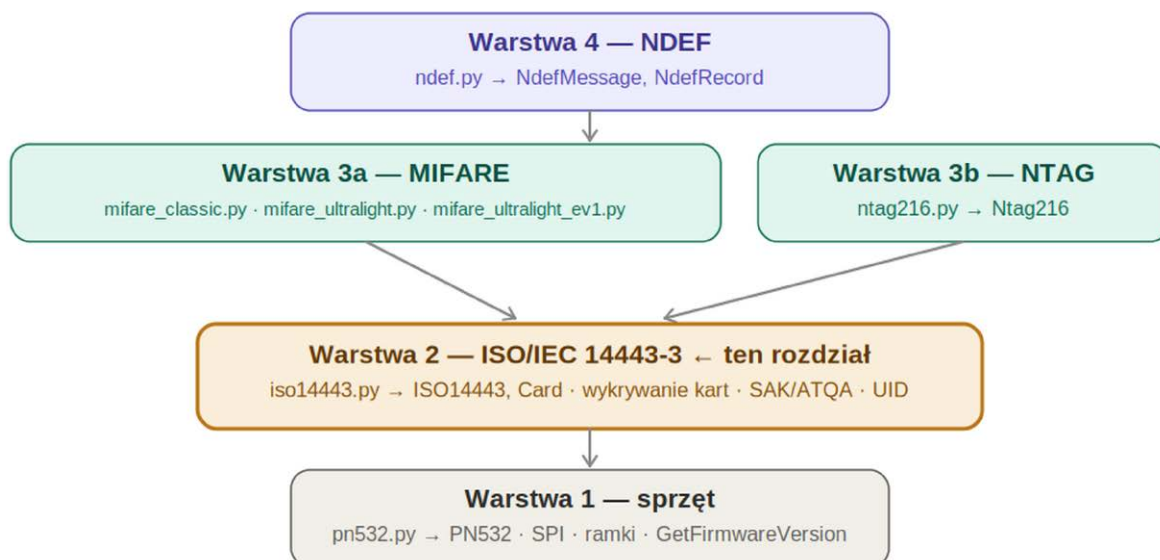
W rozdziale pierwszym nauczyliśmy się rozmawiać z układem PN532 na poziomie ramek SPI. Wiemy już, jak wysłać komendę i odebrać odpowiedź. Czas jednak wyjść poza sam czytnik i zająć się tym, co pojawia się w jego polu – kartami. Żeby karta w ogóle odpowiedziała na nasze pytania, musi przez nią przejść ściśle określona procedura opisana w normie ISO/IEC 14443-3. W tym rozdziale przerobimy ją krok po kroku, a na końcu napiszemy warstwę 2 naszej architektury – klasę ISO14443.

Czym jest ISO/IEC 14443-3 i dlaczego to ważne

ISO/IEC 14443 to rodzina czterech norm opisujących komunikację zbliżeniową na 13,56 MHz. Poszczególne części dzielą odpowiedzialność:

Część normy	Co opisuje
ISO 14443-1	Właściwości fizyczne kart (rozmiar, wytrzymałość)
ISO 14443-2	Zasilanie indukcyjne, modulacja ASK, podnośna
ISO 14443-3	Inicjalizacja i procedura antykolizyjna – jak wykryć kartę i odróżnić ją od innych
ISO 14443-4	Protokół transmisji wyższego poziomu (używany przez DESFire, MIFARE Plus SL3)

Nas interesuje przede wszystkim część trzecia. To ona opisuje, co czytnik musi zrobić, żeby w ogóle



1. Architektura warstwowa kodu kursu RFID. Warstwa 2 – ISO14443 (wyróżniona) – jest dodawana w tym rozdziale. Każda warstwa zna tylko warstwę bezpośrednio nożej

nawiązać kontakt z kartą – zanim jeszcze przejdziemy do czytania danych czy autentykacji. Bez tej procedury karta nie odpowie na żadne polecenie.

Norma definiuje dwa typy kart: Typ A (modulacja ASK 100%, kodowanie Millera) i Typ B (modulacja ASK 10%, kodowanie NRZ). Wszystkie karty, którymi zajmujemy się w kursie – MIFARE Classic, Ultralight, NTAG – są Typu A. PN532 obsługuje oba typy, ale `InListPassiveTarget` z parametrem `0x00` szuka właśnie Typu A.

Procedura wykrywania karty krok po kroku

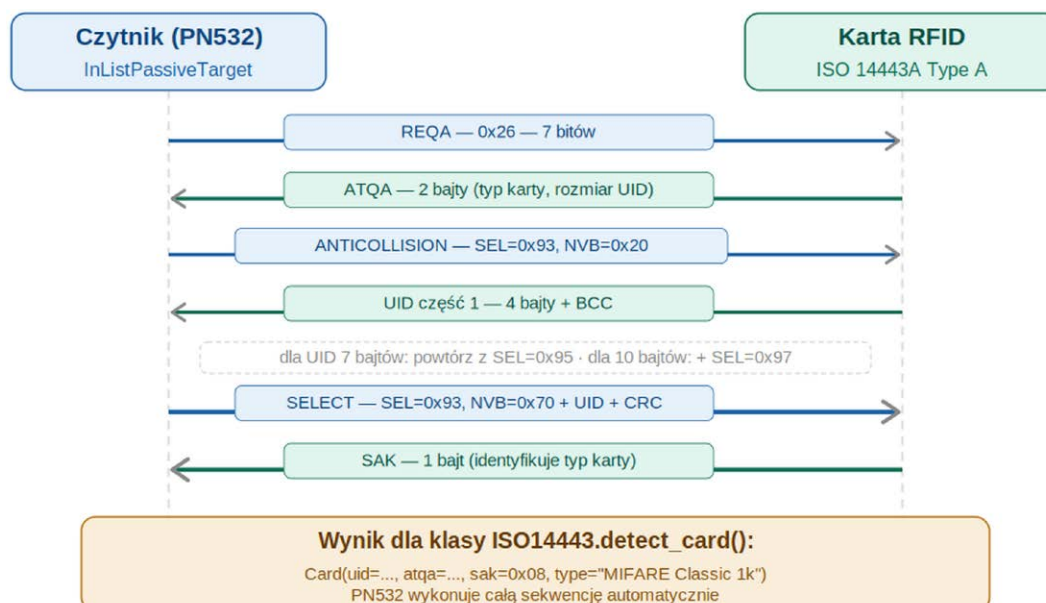
Kiedy wywołujesz `InListPassiveTarget`, PN532 samodzielnie wykonuje całą sekwencję opisaną w normie. Warto jednak zrozumieć, co się w środku dzieje – bo ATQA

i SAK, które dostajemy w odpowiedzi, są bezpośrednim efektem tej procedury.

REQA – pierwsze słowo do karty

Wszystko zaczyna się od REQA (REQuest type A) – to jednobajtowe polecenie `0x26`, nadawane ze specjalną długością 7 bitów (nie 8, jak standard SPI). Jest to jedyne polecenie w ISO 14443, które karta może odebrać bez wcześniejszej inicjalizacji. Nadawane jest do wszystkich kart w polu – to jak okrzyk ‘Czy jest tu ktoś?’.

Istnieje też wariant WUPA (`0x52`), który budzi karty będące w stanie uśpienia (HALT). W naszej uproszczonej implementacji używamy tylko REQA, bo PN532 robi to sam. Karty w trybie HALT nie odpowiadają na REQA, ale na WUPA – to ważne przy wielokrotnym odczycie tej samej karty.



2. Sekwencja inicjalizacji ISO/IEC 14443-3 Type A: REQA → ATQA → ANTICOLLISION → UID → SELECT → SAK. PN532 wykonuje ją automatycznie przez `InListPassiveTarget`

ATQA – pierwsza odpowiedź karty

Karta, która usłyszy REQA i jeszcze nie jest zaznaczona (nie przeszła SELECT), odpowiada dwubajtowym słowem ATQA (Answer To Request Type A). Z ATQA możemy odczytać kilka informacji:

Bity ATQA	Znaczenie
Bity 0-3 (nibble, inaczej półbajt, dolny bajtu 1)	Rozmiar UID: 00=4 bajty, 01=7 bajtów, 10=10 bajtów
Bit 6 bajtu 1	1 = karta wspiera antykolizję (zawsze 1 dla MIFARE)
Bity 8-15 (bajt 2)	Informacje producenta, typ pamięci – interpretacja zależna od producenta

W praktyce nie musisz analizować ATQA do identyfikacji rodziny kart – do tego służy SAK, który dostajemy po chwili. ATQA przydaje się głównie do wstępnej oceny, czy w polu jest cokolwiek, i do wykrywania specyficznych kart (np. chińskie backdoorowe karty mają ATQA=0x0004 i SAK=0x88).

Pętla antykolizyjna – dlaczego w ogóle istnieje

Wyobraź sobie, że w polu czytnika leżą dwie karty jednocześnie. Obie usłyszą REQA i obie chcą odpowiedzieć ATQA. Gdyby odpowiedziały w tym samym momencie, sygnały by się nałożyły i czytnik nie odczytałby żadnej. Pętla antykolizyjna rozwiązuje ten problem.

Algorytm jest opisany w ISO 14443-3 sekcja 6: czytnik stopniowo zawęża okno selekcji na podstawie bitów UID. Karta, której UID zaczyna się od podanego prefiksu, odpowiada; pozostałe milczą. Po kilku rundach czytnik identyfikuje dokładnie jedną kartę i ją ‘zaznacza’ komendą SELECT.

Uproszczenie: obsługujemy tylko jedną kartę

W kursie celowo upraszczamy pętlę antykolizyjną do obsługi jednej karty naraz. InListPassiveTarget z parametrem MaxTg=1 zatrzymuje się na pierwszej znalezionej karcie i nie oddaje pola innym. To wystarczy do wszystkich ćwiczeń w kursie. Pełna obsługa wielu kart jednocześnie wymagałaby iterowania przez kolejne zaznaczenia i zarządzania numerami Tg (target number). PN532 obsługuje do 2 kart jednocześnie – ale to temat na osobny artykuł.

SELECT i SAK – finał procedury

Po antykolizji czytnik wysyła SELECT z pełnym UID karty. Karta potwierdza komendą SAK (Select Acknowledge). Bajt SAK niesie kluczową informację: czy UID jest kompletny (karta 4-bajtowa) czy wymaga kolejnego poziomu kaskadowania (karty 7 i 10-bajtowe), a przede wszystkim – z jakim typem karty mamy do czynienia.

UID – 4, 7 lub 10 bajtów

Identyfikator UID może mieć trzy długości, zależnie od rodziny kart:

- 4-bajtowy UID: starsze karty MIFARE Classic. Ostatnio coraz rzadziej spotykane. Przechodzi przez jeden poziom SELECT (z SEL=0x93).
- 7-bajtowy UID: MIFARE Ultralight, NTAG, nowsze MIFARE Classic. Dwa poziomy SELECT (0x93, potem 0x95). Pierwsze 4 bajty nadawane z flagą kaskadowania (CT=0x88 jako bajt 0 w pierwszym poziomie).
- 10-bajtowy UID: rzadkie karty specjalne. Trzy poziomy SELECT (0x93 → 0x95 → 0x97).

PN532 obsługuje wszystkie trzy długości automatycznie w InListPassiveTarget. W odpowiedzi dostajemy kompletny UID bez bajtów CT – czyli czyste 4, 7 lub 10 bajtów identyfikujące kartę.

Identyfikacja rodziny kart na podstawie SAK

To najbardziej praktyczna sekcja tej teorii. Bajt SAK wystarczy, żeby zanim jeszcze dotkniemy pamięci karty, wiedzieć, czy mamy Classic, Ultralight, NTAG, czy coś zupełnie innego. Poniższa tabela pochodzi z noty aplikacyjnej NXP AN10833 ‘MIFARE type identification procedure’:

Typ karty	ATQA (hex)	SAK (hex)	Uwagi
MIFARE Classic 1k	0x0004	0x08	najpopularniejsza karta MIFARE
MIFARE Classic 4k	0x0002	0x18	
MIFARE Classic 1k (Magic/CN)	0x0004	0x88	chińska karta z backdoorem, UID zmieniany
MIFARE Ultralight	0x0044	0x00	i NTAG21x – nierozróżnialne po SAK!
MIFARE Ultralight EV1	0x0044	0x00	jak Ultralight
NTAG213/215/216	0x0044	0x00	jak Ultralight – odróżniamy po wersji bloku
MIFARE DESFire EV1/EV2	0x0344	0x20	ISO 14443-4, potrzeba RATS
MIFARE Plus SL1	0x0004	0x08	wygląda jak Classic 1k!
MIFARE Plus SL3	0x0004	0x20	ISO 14443-4

NTAG a Ultralight mają identyczne SAK i ATQA!

To częste źródło nieporozumień. Po samym SAK i ATQA nie odróżnisz NTAG216 od Ultralight EV1. Żeby stwierdzić dokładny typ, musisz po wykryciu karty odczytać jej blok konfiguracyjny (strona wersji) poleceniem GET_VERSION – to zrobimy w rozdziałach 5...7.

Na potrzeby rozdziału 2 wystarczy nam rozróżnienie: Classic/Ultralight+NTAG/ISO 14443-4.

Chińskie karty z backdoorem – jak je wykryć

Na rynku od lat dostępne są chińskie klony kart MIFARE Classic, które na pozór wyglądają identycznie jak oryginały NXP. Różnią się jednak jedną kluczową cechą: mają backdoor pozwalający na zmianę zawartości bloku 0 – czyli sektora, w którym zapisany jest m.in. UID.

W oryginalnych kartach NXP blok 0 jest zaprogramowany fabrycznie i całkowicie zablokowany przed zapisem. Zmiana UID jest niemożliwa. W chińskich klonach producent dodał ukryte polecenie, które otwiera ten blok. Karty te rozpoznaje się po SAK=0x88 (zamiast standardowego 0x08 dla Classic 1k) albo po odpowiedzi na polecenie 0x40, które w normie ISO 14443-3 nie istnieje.

Po co te karty w kursie?

Chińskie karty z backdoorem są użyteczne do ćwiczeń: możesz zmieniać ich UID, co pozwala testować kod bez ryzyka uszkodzenia oryginalnej karty. W rozdziałach 3 i 4 pokażemy, jak korzystać z tego backdoora programowo.

W zestawie do kursu (sklep.avt.pl) znajdziesz zarówno oryginalne karty NXP, jak i karty Gen1a z backdoorem. Na opakowaniu lub w opisie będą oznaczone jako 'Magic Card' lub 'UID changeable'.

Metoda detekcji w kodzie jest prosta: wysyłamy do karty bajt 0x40 przez InDataExchange. Prawdziwa karta NXP odrzuca to polecenie błędem (lub milczeniem).

Karta z backdoorem odpowiada 0x0A (ACK) – to sygnał, że mamy do czynienia z klonem.

Piszemy klasę ISO14443

Mając teorię za sobą, czas na kod. Tym razem piszemy plik iso14443.py – to warstwa 2 naszej architektury. Klasa ISO14443 przyjmuje obiekt PN532 (warstwę 1) i oferuje publiczny interfejs: detect_card, wait_for_card i detect_backdoor_card.

Zwróć uwagę na kilka projektowych decyzji, które warto zapamiętać na resztę kursu. Po pierwsze, ISO14443 nigdy nie tworzy sam obiektu PN532 – dostaje go przez konstruktor. Dzięki temu można łatwo podmienić warstwę 1 bez ruszania klasy ISO. Po drugie, detect_card zwraca None zamiast wyjątek gdy karty nie ma – to wygodniejsze w pętlach polling. Po trzecie, Card to osobna klasa danych, nie słownik – to ułatwia późniejsze rozszerzanie.

Listing 4: iso14443.py – klasa ISO14443 (warstwa 2)

```
# iso14443.py
# Warstwa 2: procedura wykrywania i identyfikacji kart ISO/IEC 14443-3 Type A
# Kurs RFID, rozdział 2

import time
from pn532 import PN532 # warstwa 1 z odcinka 1

# -----
# Komendy PN532 używane w tej warstwie
# (źródło: PN532 User Manual UM10232, rozdz. 7)
# -----

CMD_RF_CONFIGURATION      = 0x32 # konfiguracja toru RF
CMD_IN_LIST_PASSIVE_TARGET = 0x4A # wyszukaj kartę ISO 14443A

# Maksymalna liczba kart, o którą pytamy w jednym wywołaniu
_MAX_TARGETS = 1

# Typ karty: ISO 14443A (MIFARE, NTAG itp.)
_BAUD_ISO14443A = 0x00

# Typy kart na podstawie SAK – tabela z NXP AN10833
CARD_UNKNOWN      = 'Nieznany'
CARD_MIFARE_CLASSIC_1K = 'MIFARE Classic 1k'
CARD_MIFARE_CLASSIC_4K = 'MIFARE Classic 4k'
CARD_MIFARE_ULTRALIGHT = 'MIFARE Ultralight/NTAG'
CARD_ISO14443_4      = 'ISO 14443-4 (np. DESFire)'

class ISO14443Error(Exception):
    """Wyjątek warstwy ISO 14443-3."""
    pass

class Card:
    """Reprezentuje kartę wykrytą w polu czytnika."""

    def __init__(self, uid, atqa, sak):
        self.uid = bytes(uid) # unikalny identyfikator karty
        self.atqa = atqa      # Answer To Request Type A (2 bajty)
        self.sak = sak        # Select Acknowledge (1 bajt)
        self.type = _identify_card(sak)

    def uid_hex(self):
        """Zwróć UID jako czytelny ciąg hex, np. '04:A3:1B:7C'."""
        return ':'.join(f'{b:02X}' for b in self.uid)
```



```

def __repr__(self):
    return (f'Card(uid={self.uid_hex()}, '
            f'atqa={self.atqa:#06x}, sak={self.sak:#04x}, '
            f'type={self.type!r})')

def _identify_card(sak):
    """Rozpoznaj typ karty na podstawie bajtu SAK.
    Źródło: NXP AN10833 'MIFARE type identification procedure'.
    """
    # Maska bitu 6 (ISO 14443-4 compliant) i bitu 5 (UID not complete)
    sak = sak & 0xFF
    if sak & 0x04: # bit 2 ustawiony: UID niekompletny (cascade level)
        return CARD_UNKNOWN # nie powinno się zdarzyć przy normalnym odczycie
    if sak == 0x08 or sak == 0x88:
        return CARD_MIFARE_CLASSIC_1K
    if sak == 0x18:
        return CARD_MIFARE_CLASSIC_4K
    if sak == 0x00:
        return CARD_MIFARE_ULTRALIGHT # Ultralight, NTAG, Ultralight EV1
    if sak & 0x20:
        return CARD_ISO14443_4
    return CARD_UNKNOWN

class ISO14443:
    """Warstwa 2: obsługa protokołu ISO/IEC 14443-3 Type A.

    Używa klasy PN532 (warstwa 1) do wysyłania komend.
    Implementuje uproszczoną pętlę antykolizyjną dla jednej karty.
    """

    def __init__(self, pn532: PN532):
        self._pn532 = pn532
        self._rf_on = False

    def rf_on(self):
        """Włącz pole RF czytnika."""
        if not self._rf_on:
            self._pn532._send_command(
                CMD_RF_CONFIGURATION,
                bytes([0x01, 0x01]) # sterowanie RF, włącz
            )
            self._rf_on = True
            time.sleep_ms(50) # chwila na ustabilizowanie pola

    def rf_off(self):
        """Wyłącz pole RF (oszczędność energii, unikanie interferencji)."""
        if self._rf_on:
            self._pn532._send_command(
                CMD_RF_CONFIGURATION,
                bytes([0x01, 0x00]) # sterowanie RF, wyłącz
            )
            self._rf_on = False

    def detect_card(self, timeout_ms=500):
        """Wykryj jedną kartę ISO 14443A w polu czytnika.

        Wysyła komendę InListPassiveTarget do PN532, która wewnętrznie
        wykonuje całą sekwencję: REQA → antykolizja → SELECT.
        Jeśli karta zostanie znaleziona, zwraca obiekt Card.
        Jeśli nie - zwraca None.
        """
        self.rf_on()
        try:
            self._pn532._send_command(
                CMD_IN_LIST_PASSIVE_TARGET,
                bytes([_MAX_TARGETS, _BAUD_ISO14443A])
            )
            # Odpowiedź: [NbTg, Tg, ATQA(2), SAK(1), NFCIDLen(1), NFCID...]
            # Maksymalny UID to 10 bajtów (triple cascade)
            raw = self._pn532._read_response(
                CMD_IN_LIST_PASSIVE_TARGET,
                data_length=17
            )
        except Exception:
            return None

        nb_tg = raw[0] # liczba znalezionych kart
        if nb_tg == 0:

```

```

        return None

    # Parsuj odpowiedź dla pierwszej karty (Tg=1)
    # raw[1] = numer targetu (zawsze 1 przy _MAX_TARGETS=1)
    atqa = (raw[3] << 8) | raw[2] # ATQA: 2 bajty, LSB first
    sak = raw[4] # SAK: 1 bajt
    uid_len = raw[5] # długość UID w bajtach
    uid = raw[6:6 + uid_len] # sam UID

    return Card(uid=uid, atqa=atqa, sak=sak)

def wait_for_card(self, timeout_ms=5000, poll_ms=200):
    """Czekaj na kartę, odpytując co poll_ms milisekund.

    Zwraca obiekt Card gdy karta zostanie znaleziona,
    lub None po upływie timeout_ms.
    """
    t0 = time.time()
    while True:
        card = self.detect_card()
        if card is not None:
            return card
        if time.time() - t0 > timeout_ms:
            return None
        time.sleep(poll_ms)

def detect_backdoor_card(self):
    """Sprawdź, czy karta obsługuje chińskie polecenie backdoor.

    Chińskie klony MIFARE Classic z możliwością zmiany UID
    odpowiadają na polecenie 0x40 (spoza normy ISO 14443-3).
    Prawdziwe karty NXP nie odpowiadają na to polecenie.
    Zwraca True jeśli karta ma backdoor, False w przeciwnym wypadku.
    """
    # Wyślij magiczne polecenie bezpośrednio – poza normalnym API
    # InDataExchange: wyślij 1 bajt 0x40 do karty o numerze Tg=1
    CMD_IN_DATA_EXCHANGE = 0x40
    try:
        self._pn532._send_command(
            CMD_IN_DATA_EXCHANGE,
            bytes([0x01, 0x40]) # Tg=1, dane=0x40
        )
        resp = self._pn532._read_response(CMD_IN_DATA_EXCHANGE, data_length=2)
        # Karta z backdoorem odpowie 0x0A (ACK) zamiast błędu
        return resp[0] == 0x00 and resp[1] == 0x0A
    except Exception:
        return False

```

Listing 5: main2.py – skaner kart

```

# main2.py
# Rozdział 2: skanowanie kart i identyfikacja typów

from machine import SPI, Pin
from pn532 import PN532
from iso14443 import ISO14443

# --- Inicjalizacja warstwy 1 (identyczna jak w odcinku 1) ---
spi = SPI(0, baudrate=1_000_000, polarity=0, phase=0,
          sck=Pin(18), mosi=Pin(19), miso=Pin(16))
czytnik_hw = PN532(spi=spi, cs_pin=17)

# --- Inicjalizacja warstwy 2 ---
iso = ISO14443(czytnik_hw)

print('Skaner kart RFID – rozdział 2')
print('Przyłóż kartę do czytnika...')
print()

while True:
    karta = iso.wait_for_card(timeout_ms=10_000)
    if karta is None:
        print('Brak karty – sprawdź połączenie lub przyłóż kartę.')
        continue

    print(f'--- Znaleziono kartę ---')
    print(f'  UID: {karta.uid_hex()} ({len(karta.uid)} bajty)')

```

```
print(f' ATQA: {karta.atqa:#06x} ({karta.atqa:016b} binarnie)')
print(f' SAK: {karta.sak:#04x} ({karta.sak:08b} binarnie)')
print(f' Typ: {karta.type}')

# Sprawdź, czy to chińska karta z backdoorem
if iso.detect_backdoor_card():
    print(' [!] Wykryto backdoor - karta z możliwością zmiany UID!')
else:
    print(' Brak backdoora (oryginał lub karta bez furtki)')
print()
```

Listing jest krótki, bo warstwa 2 robi całą robotę. Pętla main po prostu czeka na kartę i wypisuje jej dane. Jeśli chcesz przetestować detekcję backdoora, potrzebujesz chińskiej karty Magic – przyłóż normalną Classic i Magic, żeby porównać wyniki.

Wynik dla MIFARE Classic 1k (oryginał)

```
--- Znaleziono kartę ---
UID: 04:A3:1B:7C:22:3F:80 (7 bajtów)
ATQA: 0x0044 (000000000001000100 binarnie)
SAK: 0x20 (00100000 binarnie)
Typ: MIFARE Classic 1k
Brak backdoora (oryginał lub karta bez furtki)
```

Wynik dla chińskiej karty Magic (SAK=0x88)

```
--- Znaleziono kartę ---
UID: DE:AD:BE:EF (4 bajty)
ATQA: 0x0004 (00000000000000100 binarnie)
SAK: 0x88 (10001000 binarnie)
Typ: MIFARE Classic 1k
[!] Wykryto backdoor - karta z możliwością zmiany UID!
```

Pomocniczy listing: tabela SAK/ATQA dla popularnych kart

Dla wygody – ten listing warto trzymać otwartym podczas ćwiczeń. Pokazuje wartości ATQA i SAK, jakich możesz się spodziewać dla kart w Twoim zestawie testowym:

Test z różnymi kartami

Mając kilka kart z zestawu, możesz sprawdzić, czy identyfikacja działa poprawnie. Dla każdej karty program wypisze UID, ATQA, SAK i interpretację typu. Porównaj z tabelą niżej:

```
# Przykładowe wartości ATQA i SAK dla popularnych kart
# (źródło: NXP AN10833 MIFARE type identification procedure)

# ATQA to słowo 2-bajtowe zwracane przez kartę na REQA
# SAK to 1 bajt potwierdzenia po komendzie SELECT

# Karta          ATQA (hex)  SAK (hex)  Uwagi
# -----
# MIFARE Classic 1k      0x0004      0x08      najpopularniejsza
# MIFARE Classic 4k      0x0002      0x18
# MIFARE Classic 1k (C.M) 0x0004      0x88      chińska 'magic card'
# MIFARE Ultralight      0x0044      0x00      i NTAG21x
# MIFARE Ultralight EV1   0x0044      0x00      jak Ultralight
# NTAG213/215/216        0x0044      0x00      jak Ultralight
# MIFARE DESFire EV1      0x0344      0x20      ISO 14443-4
# MIFARE Plus (SL1)       0x0004      0x08      jak Classic 1k!
# MIFARE Plus (SL3)       0x0004      0x20      ISO 14443-4
```

Uruchomienie i weryfikacja Wgrywanie pliku

Zapisz iso14443.py na Pico (File → Save As → Raspberry Pi Pico → iso14443.py), tak samo jak pn532.py w rozdziale 1. Plik pn532.py musi już być na Pico. Następnie wgryj main2.py i uruchom przez F5.

- Na Pico powinny być teraz trzy pliki: pn532.py (warstwa 1), iso14443.py (warstwa 2), main2.py (program testowy).
- Sprawdź w Thonny View → Files, że widzisz je po stronie Raspberry Pi Pico, nie tylko na lokalnym dysku.

Karta do przetestowania	Oczekiwany SAK
MIFARE Classic 1k (biała karta)	0x08 lub 0x20 (MIFARE Plus SL1)
MIFARE Classic 4k	0x18
Chińska Magic Card (UID changeable)	0x88 + wykrycie backdoora
MIFARE Ultralight (karta lub brelok)	0x00
NTAG216 (tagi NFC)	0x00 – jak Ultralight
Karta bankomatowa lub karta miejska	0x20 (ISO 14443-4, DESFire lub podobna)

Moja karta bankomatowa/bilet miesięczny nie daje MIFARE Classic?

To normalne. Nowoczesne karty płatnicze i bilety komunikacji miejskiej używają MIFARE DESFire lub ISO 14443-4 – stąd SAK=0x20. Nasz kod poprawnie identyfikuje je jako 'ISO 14443-4' i nie próbuje czytać danych MIFARE. To dobrze – DESFire ma szyfrowanie i kurs się nim nie zajmuje.

Materiały do dalszej lektury

ISO/IEC 14443-3:2018, sekcja 6 – pełna specyfikacja procedury inicjalizacji i antykolizji Typ A (tekst normy lub opracowania online).

NXP AN10833 'MIFARE type identification procedure' – tabela SAK/ATQA dla wszystkich rodzin kart NXP, dostępna bezpłatnie na nxp.com.

PN532 User Manual (UM10232), rozdz. 7.3.5 'InListPassive-Target' – opis komendy, format odpowiedzi, obsługa UID 4/7/10-bajtowych.

Podsumowanie

Ten rozdział był gęsty od teorii, ale efekty są namacalne. Mamy teraz działający skaner kart, który potrafi wykryć kartę, wyciągnąć z niej UID, ATQA i SAK, rozpoznać rodzinę karty i ostrzec przed chińskim backdoorem. Kod jest zorganizowany w dwie warstwy: PN532 zajmuje się SPI i protokołem czytnika, ISO14443 zajmuje się logiką wykrywania kart.

Warto się trzymać architektury warstwowej – w kolejnych rozdziałach każda rodzina kart dostanie własną klasę

(MifareClassic, MifareUltralight itd.), a wszystkie będą korzystać z ISO14443 przez publiczny interfejs `detect_card`. Klasy te nie będą wiedziały nic o PN532 ani o SPI – i właśnie o to chodzi.

W rozdziale 3 schodzimy głębiej do MIFARE Classic: mapa pamięci, sektory, bloki, klucze A/B i pierwsze operacje zapisu i odczytu danych.

Kurs RFID z PN532, rozdział 3

MIFARE Classic: mapa pamięci, autentykacja, zapis i odczyt

W poprzednich rozdziałach nauczyliśmy się wykrywać kartę i odczytywać jej typ na podstawie SAK i ATQA. Wiemy już, że MIFARE Classic 1k odpowiada SAK=0x08 i że przechowuje dane w 64 blokach. Czas zejść o poziom głębiej – do samej pamięci karty. W tym rozdziale napiszemy klasę `MifareClassic`, która pozwoli czytać i pisać dane w dowolnym bloku, a przy okazji poznamy mechanizm autentykacji `Crypto1` i dowiemy się, dlaczego bity dostępu są tak ważne.

Mapa pamięci MIFARE Classic

MIFARE Classic to rodzina kart NXP działająca na 13,56 MHz według normy ISO/IEC 14443-3 Type A. Istnieją dwa warianty: Classic 1k z pojemnością 1 kilobajta i Classic 4k z pojemnością 4 kilobajtów. Mimo że nazwa sugeruje pojemność, obie karty mają tę samą architekturę pamięci – różnią się tylko liczbą sektorów.

Classic 1k: 16 sektorów, 64 bloki

Pamięć Classic 1k jest podzielona na 16 sektorów (numerowanych 0...15). Każdy sektor zawiera dokładnie 4 bloki po 16 bajtów. Daje to 64 bloki i 1024 bajty łącznie. Bloki są numerowane od 0 do 63, co ważne przy adresowaniu komend `READ` i `WRITE`.

Nie wszystkie bloki są dostępne do swobodnego zapisu. Karta ma dwa rodzaje bloków o specjalnym przeznaczeniu:

- Blok 0 (*Manufacturing Block*) – zawiera 4-bajtowy UID karty, bajt BCC (checksum UID), bajt SAK i dane producenta. Jest zapisany fabrycznie i chroniony przed nadpisaniem przez sprzęt karty. W chińskich kartach Magic Gen1a ten blok jest odblokowywany specjalnym poleceniem `backdoor` – to właśnie ten mechanizm omawialiśmy w rozdziale 2.
- Blok 3 każdego sektora (*Sector Trailer*) – zawiera klucze kryptograficzne A i B oraz bity dostępu. To z tym blokiem bądźmy ostrożni: błędny zapis bitów dostępu może trwale zablokować sektor.

Sektor	Blok	Zawartość (16 bajtów)	Abs. nr
0	0	UID (4B) + BCC + SAK + dane producenta	0
	1	blok danych	1
	2	blok danych	2
	3	Klucz A (6B) Bity dostępu (4B) Klucz B (6B)	3
1	0-2	bloki danych (bloki 4-6)	4-6
	3	Sector Trailer (blok 7)	7
2	0-2	bloki danych (bloki 8-10)	8-10
	3	Sector Trailer (blok 11)	11
3-14			
15	0-2	bloki danych (bloki 60-62)	60-62
	3	Sector Trailer (blok 63)	63
12-59			

■ Blok 0 — Manufacturing Block (UID, tylko odczyt)

■ Bloki danych (wolne do zapisu po autentykacji)

■ Sector Trailer (klucze A i B, bity dostępu)

Classic 1k: 16 sektorów × 4 bloki = 64 bloki × 16 B = 1024 B

Classic 4k: sektory 0-31 po 4 bloki + sektory 32-39 po 16 bloków = 4096 B

1. Mapa pamięci MIFARE Classic 1k: 16 sektorów × 4 bloki = 64 bloki. Czerwony = Manufacturing Block (UID, tylko odczyt), zielony = bloki danych, żółty = Sector Trailer (klucze + bity dostępu). Źródło: opracowanie własne na podstawie: NXP MF1S503x MIFARE Classic EV1 1K Datasheet, Rev. 3.2 – nxp.com

Classic 4k: 40 sektorów, 256 bloków

Classic 4k dodaje więcej sektorów przy zachowaniu tej samej architektury. Sektory 0...31 działają identycznie jak w Classic 1k – po 4 bloki każdy. Sektory 32...39 są większe i zawierają po 16 bloków. Łącznie daje to 256 bloków i 4096 bajtów. Klasa MifareClassic obsługuje oba

Jak sprawdzić, czy masz Classic 1k czy 4k?

Po SAK: Classic 1k ma SAK=0x08 (lub 0x88 dla kart Magic), Classic 4k ma SAK=0x18. Kod z rozdziału 2 wypisuje te wartości dla każdej przyłożonej karty.

Po ATQA: Classic 1k → 0x0004, Classic 4k → 0x0002. Oba warianty wymagają tych samych komend READ i WRITE – różnią się tylko zakresem numerów bloków.

warianty automatycznie – metoda `block_to_sector()` zwraca prawidłowy numer sektora dla każdej długości UID i numeru bloku.

Sector Trailer – klucze i bity dostępu

Każdy sektor ma dedykowany blok Sector Trailer (blok 3 w Classic 1k, blok 15 lub 63 w Classic 4k). To 16 bajtów o ściśle określonej strukturze – trzy pola, które razem kontrolują dostęp do całego sektora:

Klucz A i klucz B

Karta ma dwa niezależne klucze na sektor: A i B, każdy o długości 6 bajtów. Fabrycznie oba są ustawione na 0xFF 0xFF 0xFF 0xFF 0xFF 0xFF – to klucz fabryczny, znany każdemu. Dlatego nowa karta bez konfiguracji jest otwarta dla każdego czytnika, który zna ten klucz.

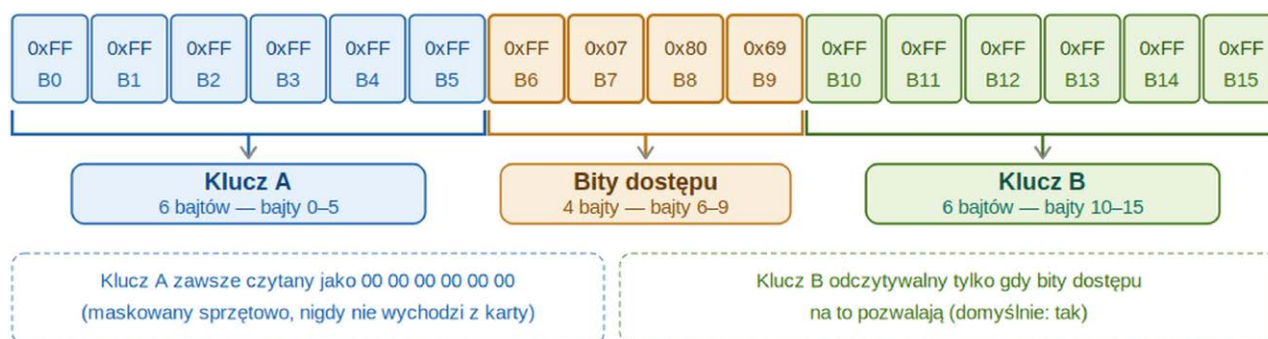
Różnica między kluczami A i B leży w bitach dostępu: domyślnie klucz A może czytać i pisać bloki danych, a klucz B może to samo lub więcej. W konfiguracji produkcyjnej zazwyczaj klucz A służy do odczytu (dystrybuowany do czytników), a klucz B do zapisu (trzymany w tajemnicy na serwerze). Po odczytaniu Sector Trailer klucz A jest zawsze maskowany zerami – nigdy nie odczytasz go z karty po skonfigurowaniu. Klucz B jest odczytywalny dopóki bity dostępu na to pozwalają.

Bity dostępu – cztery bajty, które rządzą sektorem

Bajty 6...9 Sector Trailer to bity dostępu. Kodują uprawnienia do każdego z czterech bloków sektora (trzy bloki danych + Sector Trailer) jako trzy bity C1, C2, C3. Łącznie 12 bitów jest zapisanych dwukrotnie w formie uzupełnień dla wykrycia błędów.

Poniższa tabela pokazuje najczęściej spotykane kombinacje dla bloków danych:

Sector Trailer — ostatni blok każdego sektora (16 bajtów)



Wartości fabryczne: klucze = FF FF FF FF FF FF, bity dostępu = FF 07 80 69

Fabryczny klucz 0xFF... jest znany publicznie — zmień go przed wdrożeniem produkcyjnym!

2. Format bloku Sector Trailer (16 bajtów): Klucz A (bajty 0...5, niebieski), Bity dostępu (bajty 6...9, żółty), Klucz B (bajty 10...15, zielony). Wartości fabryczne: klucze = FF FF FF FF FF FF, bity dostępu = FF 07 80 69. Źródło: opracowanie własne na podstawie: NXP MF1S503x Datasheet, sekcja 8.6 – nxp.com

C1 C2 C3	Odczyt	Zapis	Zastosowanie
0 0 0	A lub B	A lub B	domyślne – pełny dostęp
1 0 0	A lub B	B	odczyt dla wszystkich, zapis tylko B
1 1 0	A lub B	–	tylko odczyt (read-only)
0 0 1	A lub B	A lub B	blok wartości (licznik)
0 1 1	B	B	zapis i odczyt tylko kluczem B
1 1 1	–	–	trwale zablokowany!

Uwaga: błędne bity dostępu = trwałe zablokowanie sektora

Bity dostępu są zapisywane dwukrotnie (normalnie i w uzupełnieniu do 1). Karta sprawdza spójność przy każdym dostępie – niespójne bity = sektor zablokowany na zawsze.

Nigdy nie pisz bitów dostępu z głowy. Zawsze używaj kalkulatora bitów dostępu (dostępne online: 'MIFARE Classic access bits calculator') lub gotowych stałych z dokumentacji NXP UM10232.

W tym rozdziale nie modyfikujemy bitów dostępu – używamy wyłącznie fabrycznych, które mają wartość 0xFF 0x07 0x80.

Autentykacja Crypto1 – jak to działa

Zanim karta odpowie na READ lub WRITE, musi uwierzytelnić dostęp do sektora. MIFARE Classic używa do tego własnego protokołu kryptograficznego zwanego Crypto1. Nie jest to standard ISO – NXP zaprojektował go wewnętrznie w 2004 roku i przez lata trzymał algorytm w tajemnicy. W 2008 roku badacze z Radboud University odtworzyli go przez inżynierię odwrotną.

Autentykacja przebiega w trzech krokach i dlatego nazywa się 3-pass authentication:

Co robi PN532 automatycznie

Dobra wiadomość: PN532 wykonuje całą sekwencję 3-pass samodzielnie po wywołaniu komendy AUTHENTICATE (0x60 lub 0x61) przez InDataExchange. W naszym kodzie wystarczy wysłać numer bloku i klucz – resztą zajmuje się sprzęt. Jeśli autentykacja się powiedzie, PN532 utrzymuje zaszyfrowany kanał Crypto1 i wszystkie kolejne komendy READ/WRITE w tym sektorze są automatycznie szyfrowane i deszyfrowane.

Autentykacja jest ważna tylko dla jednego sektora. Dostęp do sektora 2 wymaga osobnego wywołania authenticate() z kluczem dla sektora 2 – nawet jeśli sektor 1 jest już otwarty. To celowy mechanizm bezpieczeństwa: ujawnienie klucza jednego sektora nie otwiera pozostałych.

Dlaczego Crypto1 nie jest bezpieczny?

Crypto1 to szyfr strumieniowy LFSR (Linear Feedback Shift Register) o stanie 48 bitów. Po ujawnieniu w 2008 roku okazało się, że protokół ma poważne słabości: nonce karty nie jest prawdziwie losowy (zależy od czasu od ostatniego resetu), a algorytm jest podatny na atak Nested Authentication i Dark Side Attack. Praktyczne konsekwencje: mając kilka minut i narzędzia takie jak Proxmark3 lub libnfc z biblioteką mfoc, można złamać wszystkie klucze karty bez znajomości żadnego z nich. W rozdziale 4 pokażemy, jak te ataki działają koncepcyjnie i jak testować podatność własnych kart.

Do systemów wymagających rzeczywistego bezpieczeństwa używaj MIFARE DESFire EV1/EV2 (AES-128) lub MIFARE Plus w trybie SL3.

Piszemy klasę MifareClassic

Mamy teorię – czas na warstwę 3a. Klasa MifareClassic przyjmuje obiekt ISO14443 (warstwę 2) i oferuje cztery publiczne metody: authenticate(), read_block(), write_block() i read_sector(). Plik wgrywamy na Pico obok pn532.py i iso14443.py – te trzy pliki tworzą stos warstw 1–3a.



3. Sekwencja 3-pass Crypto1 authentication MIFARE Classic. PN532 wykonuje całą sekwencję automatycznie – po wywołaniu authenticate() czytnik utrzymuje zaszyfrowany kanał Crypto1 dla całego sektora (nonce – number used once, liczba użyta raz). Źródło: opracowanie własne na podstawie: NXP AN10922 MIFARE Classic – authentication, Rev. 1.1 – nxp.com

Listing 1: mifare_classic.py

- warstwa 3a

```
# mifare_classic.py
# Warstwa 3a: obsługa kart MIFARE Classic 1k i 4k
# Kurs RFID, rozdział 3

from iso14443 import ISO14443 # warstwa 2 z rozdziału 2

# -----
# Komendy PN532 i MIFARE Classic
# -----

CMD_IN_DATA_EXCHANGE = 0x40 # PN532: wyślij dane do karty

MIFARE_AUTHENTICATE_A = 0x60 # autentykacja kluczem A
MIFARE_AUTHENTICATE_B = 0x61 # autentykacja kluczem B
MIFARE_READ = 0x30 # odczyt bloku (16 bajtów)
MIFARE_WRITE = 0xA0 # zapis bloku (16 bajtów)

BLOCK_SIZE = 16 # rozmiar bloku w bajtach (stały)
DEFAULT_KEY = b'\xff\xff\xff\xff\xff\xff' # klucz fabryczny
KEY_A = MIFARE_AUTHENTICATE_A # stała dla czytelności kodu
KEY_B = MIFARE_AUTHENTICATE_B

class MifareError(Exception):
    """Wyjątek warstwy MIFARE Classic."""
    pass

class MifareClassic:
    """Warstwa 3a: obsługa kart MIFARE Classic 1k i 4k.

    Przyjmuje obiekt ISO14443 (warstwa 2) i wykonuje operacje
    na pamięci karty. Przed każdym READ/WRITE trzeba wywołać
    authenticate() dla odpowiedniego sektora.
    """

    def __init__(self, iso: ISO14443):
        self._iso = iso
        self._pn532 = iso._pn532

    # -----
    # Metody statyczne - przeliczanie numerów bloków i sektorów
    # -----

    @staticmethod
    def block_to_sector(block_num):
        """Zwróć numer sektora dla podanego bloku (Classic 1k/4k).

        Classic 1k: sektory 0-15, po 4 bloki (bloki 0-63)
        Classic 4k: sektory 0-31, po 4 bloki (bloki 0-127)
                   sektory 32-39, po 16 bloków (bloki 128-255)
        """
        if block_num < 128:
            return block_num // 4
        return 32 + (block_num - 128) // 16

    @staticmethod
    def sector_trailer_block(sector_num):
        """Numer absolutny bloku Sector Trailer dla danego sektora."""
        if sector_num < 32:
            return sector_num * 4 + 3 # ostatni z 4 bloków
        return 128 + (sector_num - 32) * 16 + 15 # ostatni z 16

    @staticmethod
    def is_trailer(block_num):
        """Czy podany blok to Sector Trailer?"""
        if block_num < 128:
            return block_num % 4 == 3
        return (block_num - 128) % 16 == 15

    # -----
    # Autentykacja
    # -----

    def authenticate(self, block_num, key=DEFAULT_KEY,
                    key_type=KEY_A, uid=None):
```

"""Wykonaj 3-pass Crypto1 autentykację dla sektora bloku.

PN532 wykonuje cały protokół automatycznie. Po pomyślnej autentykacji wszystkie kolejne operacje READ/WRITE w tym sektorze są szyfrowane Crypto1 na poziomie sprzętu.

Args:

block_num: numer bloku (autentykuje cały sektor)
key: 6-bajtowy klucz (domyślnie klucz fabryczny)
key_type: KEY_A (0x60) lub KEY_B (0x61)
uid: UID karty; None = odczytaj automatycznie

"""

if uid is None:

card = self._iso.detect_card()

if card is None:

raise MifareError('Karta znikła z pola czytnika')

uid = card.uid

Payload InDataExchange: Tg=1, komenda auth, numer bloku,
klucz (6 B), UID karty

payload = bytes([0x01, key_type, block_num])

payload += bytes(key) + bytes(uid)

self._pn532._send_command(CMD_IN_DATA_EXCHANGE, payload)

resp = self._pn532._read_response(

CMD_IN_DATA_EXCHANGE, data_length=2)

if resp[0] != 0x00:

sector = MifareClassic.block_to_sector(block_num)

raise MifareError(

f'Autentykacja sektora {sector} nieudana '

f'(status={resp[0]:#04x}). Sprawdź klucz.'

)

Odczyt i zapis

def read_block(self, block_num):

"""Odczytaj 16 bajtów z bloku. Wymaga wcześniejszej autentykacji."""

payload = bytes([0x01, MIFARE_READ, block_num])

self._pn532._send_command(CMD_IN_DATA_EXCHANGE, payload)

resp = self._pn532._read_response(

CMD_IN_DATA_EXCHANGE, data_length=18)

if resp[0] != 0x00:

raise MifareError(

f'READ blok {block_num} nieudany (status={resp[0]:#04x})')

return bytes(resp[1:17]) # bajty 1-16: dane; 17-18: CRC

def write_block(self, block_num, data):

"""Zapisz 16 bajtów do bloku. Wymaga wcześniejszej autentykacji."""

if len(data) != BLOCK_SIZE:

raise MifareError(

f'Dane muszą mieć dokładnie {BLOCK_SIZE} bajtów, '

f'podano {len(data})')

if block_num == 0:

raise MifareError(

'Blok 0 (dane producenta + UID) jest tylko do odczytu')

if MifareClassic.is_trailer(block_num):

raise MifareError(

f'Blok {block_num} to Sector Trailer - użyj '

f'write_trailer() zamiast write_block()')

payload = bytes([0x01, MIFARE_WRITE, block_num]) + bytes(data)

self._pn532._send_command(CMD_IN_DATA_EXCHANGE, payload)

resp = self._pn532._read_response(

CMD_IN_DATA_EXCHANGE, data_length=2)

if resp[0] != 0x00:

raise MifareError(

f'WRITE blok {block_num} nieudany (status={resp[0]:#04x})')

def read_sector(self, sector_num, key=DEFAULT_KEY,

key_type=KEY_A, uid=None):

"""Autentykuj i odczytaj wszystkie bloki sektora naraz.

Zwraca listę 4 elementów (Classic 1k, sektory 0-15)

lub 16 elementów (Classic 4k, sektory 32-39).

Każdy element to bytes(16).

"""

```

if sector_num < 32:
    first = sector_num * 4
    count = 4
else:
    first = 128 + (sector_num - 32) * 16
    count = 16

self.authenticate(first, key=key, key_type=key_type, uid=uid)
return [self.read_block(first + i) for i in range(count)]

```

Zwróć uwagę na metody statyczne `block_to_sector()` i `is_trailer()`. Są statyczne, bo nie potrzebują połączenia z kartą – tylko liczą. Możesz je wywołać bez obiektu: `MifareClassic.block_to_sector(7)` zwróci 1. Przydatne przy debugowaniu bez sprzętu.

Metoda `write_block()` daje wyjątek przy próbie zapisu do bloku 0 lub Sector Trailer. To celowe zabezpieczenie

– przypadkowy zapis kluczy przez `write_block()` zamiast dedykowanej metody `write_trailer()` byłby trudny do debugowania i mógłby trwale zablokować sektor.

Listing 2: pomocniczy – tabela bitów dostępu

```

# Domyślne bity dostępu (transport configuration – fabryczne):
# Bajty 6–8: 0xFF 0x07 0x80 (zakodowane C1,C2,C3 dla 4 bloków)
# Bajt 9: 0x69 (data byte – wolny, używany przez klucz B lub dane)

# Dekodowanie bitów dostępu dla bloku danych:
#
# Blok   C1  C2  C3   Klucz A może   Klucz B może
# -----
# 0-2    0   0   0   R W (domyślnie) R W
# 0-2    0   1   0   R               R W
# 0-2    1   0   0   R               R
# 0-2    0   0   1   R W             - (blok wartości)
# ...

# Najprostsza zmiana praw dostępu: ustaw C1C2C3 = 100 dla bloku danych
# = tylko klucz B może pisać, oba mogą czytać.
# Sektor Trailer (blok 3):
# C1C2C3 = 001 (domyślne): klucz A odczytywalny kluczem B
# C1C2C3 = 011: klucz A NIEODCZYTYWALNY (produkcja)

# UWAGA: błędne bity dostępu → TRWAŁE zablokowanie sektora!
# Przed zapisem nowych bitów zawsze weryfikuj wzorem z UM10232.

```

Listing 3: main3.py – odczyt i zapis sektora

```

# main3.py – rozdział 3: odczyt i zapis MIFARE Classic

from machine import SPI, Pin
from pn532 import PN532
from iso14443 import ISO14443
from mifare_classic import MifareClassic, DEFAULT_KEY, KEY_A

# --- Inicjalizacja warstw 1 i 2 (identyczna jak w rozdziałach 1-2) ---
spi = SPI(0, baudrate=1_000_000, polarity=0, phase=0,
          sck=Pin(18), mosi=Pin(19), miso=Pin(16))
czytnik_hw = PN532(spi=spi, cs_pin=17)
iso = ISO14443(czytnik_hw)
mifare = MifareClassic(iso)

print('MIFARE Classic – rozdział 3')
print('Przyłóż kartę MIFARE Classic 1k lub 4k...')
print()

while True:
    karta = iso.wait_for_card(timeout_ms=10_000)
    if karta is None:
        print('Brak karty – sprawdź połączenie.')
        continue

    print(f'Karta: UID={karta.uid_hex()} typ={karta.type}')

```



```

try:
    # — ODCZYT sektora 1 (bloki 4-7) —————
    print()
    print('Odczyt sektora 1 (bloki 4-7) kluczem fabrycznym:')
    bloki = mifare.read_sector(
        sector_num=1, key=DEFAULT_KEY, key_type=KEY_A,
        uid=karta.uid
    )
    for i, blok in enumerate(bloki):
        nr = 4 + i
        opis = ' ← Sector Trailer' if i == 3 else ''
        print(f' Blok {nr:2d}: {blok.hex(" ")}{opis}')

    # — ZAPIS do bloku 4 —————
    print()
    dane = b'AVT kurs RFID ' # dokładnie 16 bajtów
    print(f'Zapis do bloku 4: {dane}')
    mifare.authenticate(
        4, key=DEFAULT_KEY, key_type=KEY_A, uid=karta.uid
    )
    mifare.write_block(4, dane)

    # — WERYFIKACJA —————
    mifare.authenticate(
        4, key=DEFAULT_KEY, key_type=KEY_A, uid=karta.uid
    )
    odczyt = mifare.read_block(4)
    wynik = 'OK' if odczyt == dane else 'NIEZGODNOSC!'
    print(f'Weryfikacja: {wynik} ({odczyt})')

except Exception as e:
    print(f'Błąd: {e}')

print()

```

Uruchomienie i weryfikacja

Wgraj na Pico trzy pliki: pn532.py (rozdział 1), iso14443.py (rozdział 2) i nowy mifare_classic.py. Następnie wgraj main3.py i uruchom przez F5. Przyłóż kartę MIFARE Classic 1k – powinieneś zobaczyć zawartość sektora 1 w konsoli, zapis do bloku 4 i potwierdzenie odczytu.

Przykładowy wynik dla nowej karty z kluczem fabrycznym:

Test z Classic 4k

Podmieniając kartę na Classic 4k, program działa bez zmian – metoda read_sector(1) zwróci identycznie 4 bloki. Żeby skorzystać z większej pojemności, odwołaj się do sektorów 32...39, na przykład read_sector(32). Klasa automatycznie rozpozna duży sektor i zwróci 16 bloków zamiast 4.

Oczekiwany wynik w konsoli

MIFARE Classic – rozdział 3

Przyłóż kartę MIFARE Classic 1k lub 4k...

Karta: UID=04:A3:1B:7C:22:3F:80 typ=MIFARE Classic 1k

Odczyt sektora 1 (bloki 4-7) kluczem fabrycznym:

```

Blok 4: 00 00 00 00 00 00 00 00 00 00 00 00 00 00 00 00
Blok 5: 00 00 00 00 00 00 00 00 00 00 00 00 00 00 00 00
Blok 6: 00 00 00 00 00 00 00 00 00 00 00 00 00 00 00 00
Blok 7: 00 00 00 00 00 00 ff 07 80 69 ff ff ff ff ff ff ← Sector Trailer

```

```

Zapis do bloku 4: b'AVT kurs RFID '
Weryfikacja: OK (b'AVT kurs RFID ')

```

Sektor Trailer (blok 7) zawiera same zera w miejscu kluczy A i B – to normalne, bo klucz A jest zawsze maskowany przy odczycie. Bity dostępu 0xFF 0x07 0x80 to konfiguracja fabryczna (pełny dostęp kluczem A lub B). Bajt 0x69 to data byte Sector Trailer – wolny do zapisu.

Jeśli widzisz błąd autentykacji, karta może mieć zmienne klucze. Spróbuj kluczem 0xA0 0xA1 0xA2 0xA3 0xA4 0xA5 – to drugi popularny klucz spotykany w starszych systemach. Listę typowych kluczy i narzędzia do ich odgadywania omawiamy w rozdziale 4.

Materiały do dalszej lektury

NXP MF1S503x/MF1S703x MIFARE Classic EV1 1K/4K Datasheet – pełna specyfikacja pamięci, Sector Trailer i bitów dostępu. Pobierz bezpłatnie z nxp.com.

NXP AN10922 MIFARE Classic – Exchanging data, Rev. 1.1 – opis protokołu autentykacji i formatu komend READ/WRITE.

Flavio D. Garcia et al., 'Dismantling MIFARE Classic' (2008) – praca naukowa, która ujawniła słabości Crypto1. Dostępna jako PDF na stronach Radboud University.

Podsumowanie

Mamy teraz w pełni działającą warstwę 3a – klasę MifareClassic z autentykacją, odczytem i zapisem. Kod jest krótki, bo PN532 robi ciężką pracę: 3-pass Crypto1 authentication, szyfrowanie kanału i weryfikację CRC to wszystko sprzęt. My tylko wysyłamy klucz i numer bloku.

Architektura warstw sprawdza się w praktyce: main3.py nie importuje nic z pn532.py i nic nie wie o SPI. Gdybyśmy

jutro zmienili czytnik na inny model, wystarczyłoby podmienić warstwę 1 – reszta zostaje bez zmian.

W rozdziale 4 wejdziemy głębiej w bezpieczeństwo MIFARE Classic: zobaczymy, jak działają ataki Nested i Dark Side, jak wykryć kartę z niestandardowym kluczem i jak chronić własne dane przez właściwe bity dostępu i zmianę kluczy.

Kurs RFID z PN532, rozdział 4

MIFARE Classic: Crypto1, ataki i ochrona własnych kart

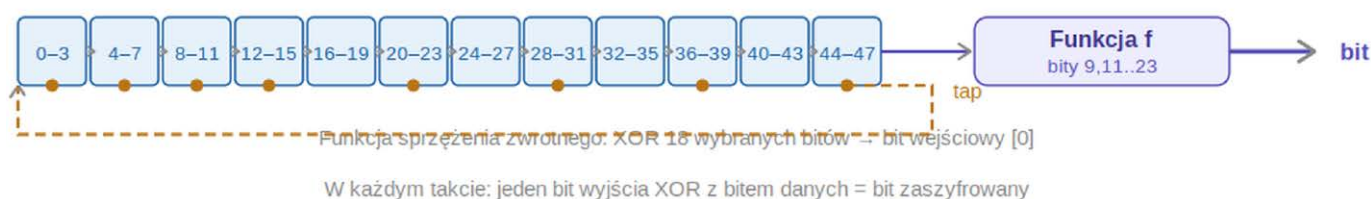
W rozdziale 3 nauczyliśmy się czytać i pisać dane na kartach MIFARE Classic. Wiemy, że autentykacja opiera się na algorytmie Crypto1 i że PN532 wykonuje ją automatycznie. W tym rozdziale wejdziemy za kulisy: zobaczymy, jak Crypto1 naprawdę działa, dlaczego jest podatny na ataki i – co ważniejsze – jak własne karty zabezpieczyć przez zmianę kluczy i właściwą konfigurację bitów dostępu.

Crypto1 – jak działa ten algorytm

Crypto1 to szyfr strumieniowy zaprojektowany przez NXP w 2004 roku specjalnie dla kart MIFARE Classic. Przez cztery lata specyfikacja algorytmu była objęta tajemnicą handlową. W 2008 roku trójka naukowców z Radboud

University w Nijmegen – Flavio Garcia, Gerhard de Koning Gans i Roel Verdult – odtworzyła go przez inżynierię odwrotną, zaglądając przez mikroskop elektronowy do warstw krzemowych karty.

Crypto1 — 48-bitowy LFSR (Linear Feedback Shift Register)



Inicjalizacja rejestru:

$LFSR[0..47] = \text{Klucz}(48b) \oplus \text{UID}(32b) \oplus \text{nonce czytnika}(32b)$

Kluczowa słabość:

nonce karty pochodzi z tego samego LFSR — przewidywalny!

- komórki LFSR (48 bitów = klucz strumieniowy)
- punkty tap — sprzężenie zwrotne (18 pozycji)
- nieliniowa funkcja wyjściowa f

Wielomian sprzężenia: $x^{48} + x^{43} + x^{39} + x^{38} + x^{36} + x^{34} + x^{33} + x^{31} + x^{29} + x^{24} + x^{23} + x^{21} + x^{19} + x^{13} + x^9 + x^7 + x^4 + x^2$ — Garcia et al. 2008

1. Schemat 48-bitowego LFSR algorytmu Crypto1. Każdy takt przesuwa bity w prawo, nowy bit wejściowy = XOR 18 wybranych pozycji (tap). Funkcja nieliniowa f generuje bit strumienia szyfrującego. Kluczowa słabość: nonce karty pochodzi z tego samego LFSR co szyfrowanie. Źródło: opracowanie własne na podstawie: Garcia et al., 'Dismantling MIFARE Classic', ESORICS 2008 – cs.ru.nl/~flaviog/publications/Dismantling.Mifare.pdf

LFSR – rejestr przesuwający ze sprzężeniem zwrotnym

Sercem Crypto1 jest 48-bitowy LFSR (*Linear Feedback Shift Register* – rejestr przesuwający ze sprzężeniem zwrotnym). Wyobraź sobie 48 przerzutników ustawionych w rząd. W każdym takcie zegara wartość każdego przerzutnika przesuwa się o jedno miejsce w prawo, a nowa wartość wpisana po lewej stronie jest funkcją XOR kilku wybranych bitów z całego rejestru. Ta prosta operacja generuje pseudolosowy strumień bitów, którym szyfrujemy dane.

Inicjalizacja kluczem i UID

Przed rozpoczęciem komunikacji LFSR jest ładowany zawartością klucza (48 bitów) przez XOR z UID karty i nonce czytelnika. Ten krok miał zapewnić, że każda sesja używa innego klucza strumieniowego – co jest dobrą zasadą kryptograficzną. Problem leży gdzie indziej: w generatorze nonce po stronie karty.

Słabość: przewidywalny generator nonce

Nonce karty (RndB, 8 bajtów) powinien być losowy – każda sesja autentykacji powinna używać innej wartości, co uniemożliwiałoby ataki powtórzeniowe. W MIFARE Classic nonce jest generowany przez ten sam LFSR co szyfrowanie – tyle że nie zasilony kluczem, lecz uruchomiony od zera po każdym przyłożeniu karty do pola. Zegar napędza LFSR z częstotliwością pola RF (13,56 MHz), więc nonce zależy bezpośrednio od czasu, który upłynął od resetu karty.

Konsekwencja: jeśli zmierzysz, ile czasu upłynęło między przyłożeniem karty a żądaniem autentykacji,

możesz przewidzieć nonce karty z dużą dokładnością. To właśnie ta słabość jest wykorzystywana przez ataki Dark Side i Nested.

Ataki na MIFARE Classic

Znamy trzy główne klasy ataków na MIFARE Classic. Każda kolejna jest bardziej wymagająca pod względem sprzętu i czasu, ale też działa w trudniejszych warunkach. Żaden z tych ataków nie jest atakiem siłowym na 48-bitowy klucz – to ataki na słabości protokołu, które redukują przestrzeń poszukiwań z 2^{48} do tysięcy lub dziesiątek tysięcy możliwości.

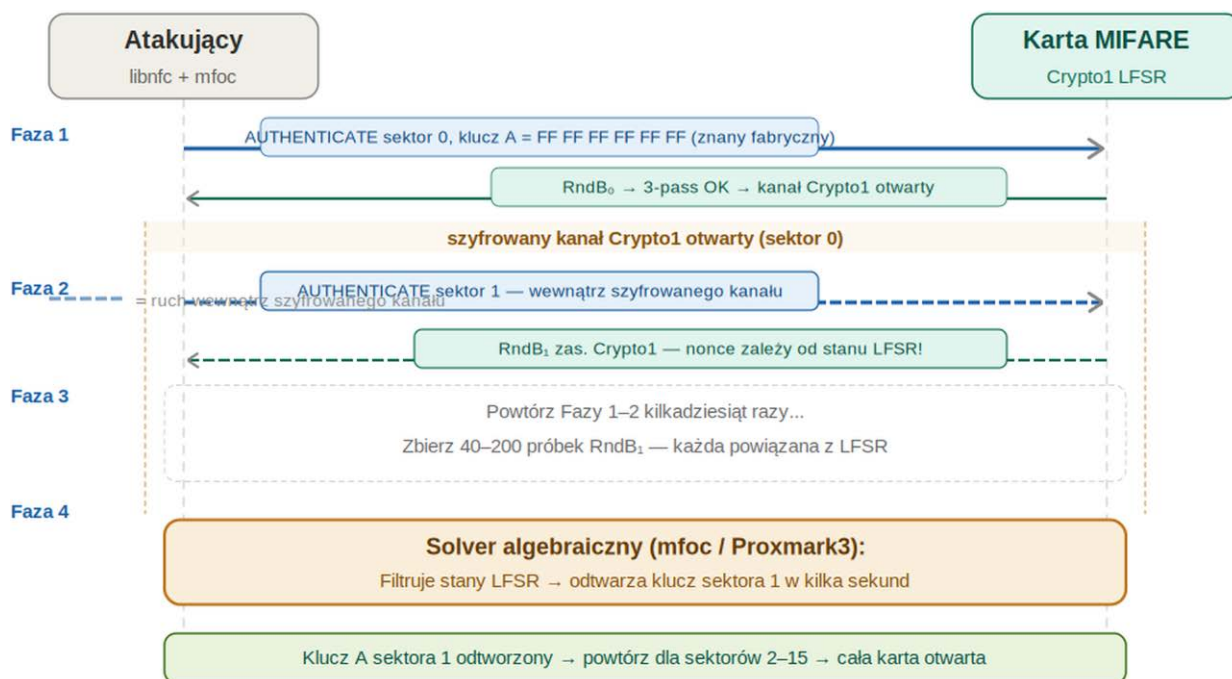
Dark Side Attack – bez znajomości żadnego klucza

Dark Side Attack (Courtois et al. 2008, implementacja: mfcuk) działa gdy nie znasz żadnego klucza na karcie. Metoda opiera się na specyficznym zachowaniu karty w sytuacji błędnej autentykacji: gdy czytnik

Jak wygląda Dark Side Attack w praktyce

1. Przytóż kartę do czytnika (Proxmark3 lub układ z libnfc + laptopem).
2. Narzędzie mfcuk wielokrotnie inicjuje autentykację, celowo podając błędny token.
3. Zbiera zaszyfrowane odpowiedzi błędu (kilkaset do kilku tysięcy).
4. Algorytm algebraiczny odtwarza stan LFSR i z niego klucz (kilka minut na laptopie).
5. Mając klucz A do jednego sektora, możesz użyć Nested Attack na pozostałe.

Czas: 5...30 minut zależnie od szczęścia z nonce. Wymagany sprzęt: Proxmark3 lub Raspberry Pi z modułem PN532 i biblioteką libnfc.



2. Sekwencja Nested Authentication Attack. Atakujący autentykuje się kluczem fabrycznym do sektora 0 (Faza 1), następnie wewnątrz zaszyfrowanego kanału żąda autentykacji do sektora 1 (Faza 2). Zbiera kilkadziesiąt nonce (Faza 3) i solver algebraiczny odtwarza klucz (Faza 4) w kilka sekund. Źródło: opracowanie własne na podstawie: de Koning Gans et al., 'A Practical Attack on the MIFARE Classic', ESORICS 2008

podaje nieprawidłowy token, karta odpowiada zaszyfrowanym komunikatem błędu. Jeśli klucz używany do szyfrowania odpowiedzi jest powiązany ze stanem LFSR w przewidywalny sposób – a tak jest w Crypto1 – można zbierać wiele takich odpowiedzi i na ich podstawie odtworzyć klucz metodą algebraiczną.

Nested Authentication Attack – znasz jeden klucz

Nested Attack (de Koning Gans et al. 2008, implementacja: mfoc) jest dużo szybszy, ale wymaga znajomości co najmniej jednego klucza – choćby fabrycznego 0xFF do dowolnego sektora. Na tym właśnie polega jego praktyczna groźba: jeśli administrator zostawił choćby jeden sektor z kluczem fabrycznym, otwiera furtkę na cały chip.

Mechanizm jest precyzyjny: po otwarciu zaszyfrowanego kanału kluczem A sektora 0 wysyłamy w środku tego samego połączenia żądanie autentykacji do sektora 1. Nonce karty dla tej drugiej autentykacji nie jest generowany świeżo – LFSR jest kontynuacją poprzedniego stanu. Zbierając kilkadziesiąt takich nonce, atakujący może użyć algebraicznego solvera do odtworzenia aktualnego stanu LFSR, a stamtąd – klucza sektora 1.

Hardnested Attack – karty z losowym nonce

NXP wprowadziło w późniejszych wariantach MIFARE Classic EV1 generowanie prawdziwie losowych nonce, co uniemożliwia Nested Attack oparty na przewidywalności. Hardnested Attack (Blapost 2015, wbudowany w Proxmark3) obchodzi tę poprawkę, zbierając znacznie więcej próbek i używając bardziej zaawansowanego filtrowania stanów LFSR. Wymaga kilku godzin i Proxmark3 – nie jest trywialny, ale działa.

Porównanie ataków

Atak	Wymagania	Czas	Narzędzia
Dark Side	Żadnego klucza	5...30 min	mfcuk + libnfc
Nested	1 klucz (choćby fabr.)	10...60 sek.	mfoc + libnfc
Hardnested	1 klucz + EV1	kilka godz.	Proxmark3
Relay	Fizyczna karta w pobliżu	chwilowy	dwa urządzenia NFC

Ważna uwaga prawna i etyczna

Opisane techniki służą do testowania własnych kart i systemów, do których masz autoryzowany dostęp. Użycie ich na cudzych kartach lub systemach jest przestępstwem. Proxmark3 i libnfc to legalne narzędzia bezpieczeństwa stosowane przez pentesterów, audytorów i badaczy.

W kursie nie dostarczamy gotowych skryptów do łamania cudzych kart – wyłącznie wyjaśniamy mechanizmy i pokazujemy, jak zabezpieczyć własne systemy przed tymi atakami.

Chińskie karty z backdoorem – Gen1a, Gen2, Gen3

W zestawie do kursu znajdziesz karty opisane jako ‘Magic Card’ lub ‘UID changeable’. To chińskie klony MIFARE Classic z modyfikacją pozwalającą na zmianę zawartości bloku 0 – bloku, który w oryginalnych kartach NXP jest trwale zablokowany po programowaniu fabrycznym. Karty te mają kilka generacji różniących się sposobem dostępu do backdoora.

Generacja	Backdoor	Wykrycie	Zastosowanie w kursie
Gen1a	Polecenie 0x40 poza normą ISO	SAK=0x88 lub resp 0x40	rozdziały 4, 9, 10
Gen2 (CUID)	Zapis do bloku 0 przez WRITE	trudne – bez spec. UID	testowanie zapisów
Gen3	GET/SET UID przez APDU	specjalne polecenie	zaawansowane testy

Gen1a – backdoor poleceniem 0x40

Gen1a (nazywane też ‘Chinese Magic Card’ lub ‘UID backdoor’) to najczęściej spotykany typ. Gdy wyślemy do karty bajt 0x40 przez InDataExchange, karta odpowiada 0x0A (ACK) zamiast odrzucić polecenie błędem. Po tej odpowiedzi możemy wysłać komendę WRITE do bloku 0 i nadpisać UID, BCC i dane producenta. Oryginalny MIFARE Classic od NXP ignoruje polecenie 0x40 całkowicie.

Wykrywanie Gen1a widziałeś już w rozdziale 2: metoda detect_backdoor_card() wysyła 0x40 i sprawdza odpowiedź. Karty Gen1a mają też SAK=0x88 zamiast standardowego 0x08 – to drugi sposób identyfikacji bez wysyłania backdoorowego polecenia.

Gen2 (CUID) – przez standardowy WRITE

Karty Gen2 są subtelniejsze: blok 0 można nadpisać zwykłą komendą WRITE po standardowej autentykacji. Nie ma specjalnego polecenia backdoor, SAK jest normalny (0x08), ATQA jest normalne (0x0004). Jediną różnicą jest to, że zapis do bloku 0 się powiedzie zamiast zwrócić błąd. Czytniki produkowane do 2015 roku zazwyczaj nie wykrywają tej modyfikacji.

Wadą Gen2 z perspektywy testowania jest brak możliwości cofnięcia zapisu. Jeśli wpiszesz nieprawidłowy UID lub BCC, karta może przestać odpowiadać na standardowe komendy. W kursie do ćwiczeń z emulacją (rozdział 9) zalecamy Gen1a – jest bezpieczniejsza w użyciu.

Jak chronić własne karty przed atakami

Dobre wiadomości: wyeliminowanie kluczy fabrycznych całkowicie usuwa możliwość Nested Attack. Dark Side Attack staje się wtedy jedyną opcją – a jest wolniejszy i bardziej wykrywalny. Poniżej cztery konkretne kroki w kolejności ważności.

Krok 1 – Zmień klucze fabryczne we wszystkich sektorach

To absolutna podstawa. Każdy sektor z kluczem 0xFF 0xFF 0xFF 0xFF 0xFF 0xFF to otwarte drzwi. Zmiana kluczy w sektorach, które rzeczywiście używasz, zajmuje sekundy – kod w tym rozdziale to robi. Sektory których nie używasz: zablokuj je bitami dostępu (C1C2C3 = 111) lub wpisz losowy klucz, którego nie przechowujesz nigdzie.

Krok 2 – Skonfiguruj bity dostępu

Domyślna konfiguracja fabryczna daje pełny dostęp zarówno kluczem A jak i B. Właściwa konfiguracja produkcyjna powinna rozdzielić uprawnienia:

- Klucz A = odczyt danych (dystrybuowany do czytników w terenie). Nawet jeśli zostanie złamany, atakujący może tylko czytać.
- Klucz B = zapis danych (przechowywany wyłącznie po stronie serwera). Wymagany do aktualizacji danych na karcie.
- Bity dostępu C1C2C3 = 100 dla bloków danych: odczyt kluczem A lub B, zapis tylko kluczem B.
- Sector Trailer: klucz A nieodczytywalny (C1C2C3 = 011 dla bloku 3). Raz wpisany klucz B pozostaje ukryty dla zewnętrznych czytników.

Krok 3 – Dywersyfikuj klucze

Jeśli wszystkie Twoje karty mają ten sam klucz A, złamanie jednej karty ujawnia klucz do wszystkich. Właściwe podejście: klucz do konkretnej karty = HMAC(klucz_główny, UID_karty). Każda karta ma inny klucz pochodny

– utrata jednej karty nie narusza bezpieczeństwa pozostałych. Implementacja wymaga serwera backendu przechowującego klucz główny.

Krok 4 – Rozważ migrację na DESFire EV2

Jeśli projektujesz nowy system i bezpieczeństwo jest krytyczne, MIFARE Classic to zły wybór – nie dlatego, że ataki są łatwe (wymagają sprzętu i wiedzy), ale dlatego, że Crypto1 jest złamanym algorytmem bez możliwości naprawy. MIFARE DESFire EV2 używa AES-128 i jest uważany za bezpieczny do zastosowań finansowych. W rozdziale 8 pokażemy, jak nawiązać sesję z DESFire przez PN532.

Kod: zmiana kluczy i walidacja bitów dostępu

Czas na implementację. Dwa nowe pliki: access_bits.py z gotowymi konfiguracjami i kalkulatorem, oraz rozszerzenie klasy MifareClassic o metodę write_trailer(). Całość demonstruje main4.py – program, który zmienia klucze w sektorze testowym i weryfikuje nową konfigurację.

Listing 1: access_bits.py – konfiguracje i walidator

Funkcja validate_access_bits() jest obowiązkowym strażnikiem przed wywołaniem write_trailer(). Sprawdza, czy trzy bajty bitów dostępu zawierają prawidłowe negacje – dokładnie tak jak sprawdza karta sprzętowo. Błędne bity powodują trwałe zablokowanie sektora, więc walidacja po stronie kodu to dodatkowy poziom ochrony.

```
# access_bits.py
# Kalkulator i walidator bitów dostępu MIFARE Classic
# Kurs RFID, rozdział 4

# Predefiniowane konfiguracje bitów dostępu (bajty 6-8 Sector Trailer)
# Bajt 9 (data byte) ustawiamy zawsze na 0x69 (wartość fabryczna).
#
# Format: (C1C2C3 dla bloku 0, bloku 1, bloku 2, Sector Trailer)
# Bity C1C2C3 dla każdego bloku zakodowane w 3 bajtach wg NXP AN10922.

# Gotowe konfiguracje - możesz używać wprost w write_trailer()

# Konfiguracja fabryczna: pełny dostęp kluczem A lub B
ACCESS_TRANSPORT = bytes([0xFF, 0x07, 0x80]) # C1C2C3 = 000 dla bloków 0-2

# Tylko odczyt: bloki 0-2 tylko do odczytu (kluczem A lub B)
ACCESS_READ_ONLY = bytes([0xFF, 0x07, 0x8F]) # C1C2C3 = 010 dla bloków 0-2

# Odczyt kluczem A, zapis tylko kluczem B
ACCESS_READ_A_WRITE_B = bytes([0xFF, 0x0F, 0x80]) # C1C2C3 = 100

def validate_access_bits(ac_bytes):
    """Sprawdź spójność 3 bajtów bitów dostępu.

    MIFARE Classic koduje każdy bit dostępu i jego negację.
    Karta sprawdza spójność przy każdym dostępie - błąd = trwałe zablokowanie.

    Returns: True jeśli bity są spójne, False jeśli nie.
    """
```

```

if len(ac_bytes) != 3:
    return False
b6, b7, b8 = ac_bytes[0], ac_bytes[1], ac_bytes[2]
# Wyciągnij bity C1, C2, C3 dla 4 bloków (12 bitów)
c1 = (b7 >> 4) & 0x0F
c2 = b8 & 0x0F
c3 = (b8 >> 4) & 0x0F
# Negacje powinny być w b6 (górny nibble) i b7 (dolny nibble)
c1_bar = (~b6 >> 4) & 0x0F
c2_bar = (~b6) & 0x0F
c3_bar = (~b7) & 0x0F
return c1 == c1_bar and c2 == c2_bar and c3 == c3_bar

def describe_access(ac_bytes):
    """Opisz słownie uprawnienia zakodowane w bitach dostępu."""
    b6, b7, b8 = ac_bytes[0], ac_bytes[1], ac_bytes[2]
    c1 = (b7 >> 4) & 0x0F
    c2 = b8 & 0x0F
    c3 = (b8 >> 4) & 0x0F
    table = {
        (0,0,0): 'R/W kluczem A lub B (transport)',
        (0,1,0): 'R kluczem A lub B, W tylko kluczem B',
        (1,0,0): 'R kluczem A lub B, W tylko kluczem B (AC mod. B)',
        (1,1,0): 'Tylko odczyt (R/O)',
        (0,0,1): 'Blok wartości - increment/decrement',
        (1,0,1): 'Blok wartości - tylko decrement',
        (0,1,1): 'Zapis i odczyt tylko kluczem B',
        (1,1,1): 'ZABLOKOWANY - brak dostępu!',
    }
    results = []
    for i in range(3): # bloki danych 0-2
        key = ((c1 >> i) & 1, (c2 >> i) & 1, (c3 >> i) & 1)
        results.append(f'Blok {i}: {table.get(key, "nieznana")}')
    return results

```

Listing 2: write_trailer() - dopisz do mifare_classic.py

```

# Rozszerzenie klasy MifareClassic - dopisz do mifare_classic.py

def write_trailer(self, sector_num, key_a, access_bits,
                  key_b, auth_key=DEFAULT_KEY, auth_key_type=KEY_A,
                  uid=None):
    """Zapisz Sector Trailer z nowymi kluczami i bitami dostępu.

    UWAGA: Operacja nieodwracalna jeśli podasz błędne bity dostępu!
    Zawsze najpierw waliduj bity przez validate_access_bits().

    Args:
        sector_num: numer sektora (0-15 dla 1k, 0-39 dla 4k)
        key_a: nowy Klucz A (6 bajtów)
        access_bits: bity dostępu (3 bajty - waliduj przed użyciem!)
        key_b: nowy Klucz B (6 bajtów)
        auth_key: aktualny klucz do autentykacji (przed zmianą)
        auth_key_type: typ aktualnego klucza (KEY_A lub KEY_B)
        uid: UID karty (None = odczytaj automatycznie)
    """
    from access_bits import validate_access_bits
    if not validate_access_bits(access_bits):
        raise MifareError(
            'Bity dostępu nie przeszły walidacji! '
            'Sprawdź wartości - błędne bity = trwałe zablokowanie.'
        )
    trailer_block = MifareClassic.sector_trailer_block(sector_num)
    self.authenticate(
        trailer_block, key=auth_key,
        key_type=auth_key_type, uid=uid
    )
    # Złóż 16 bajtów Sector Trailer
    data = bytes(key_a) + bytes(access_bits) + b'\x69' + bytes(key_b)
    # Usuń zabezpieczenie przed zapisem do traileru w write_block()
    # (wywołujemy niższy poziom - _write_raw)
    payload = bytes([0x01, MIFARE_WRITE, trailer_block]) + data
    self._pn532._send_command(CMD_IN_DATA_EXCHANGE, payload)
    resp = self._pn532._read_response(
        CMD_IN_DATA_EXCHANGE, data_length=2)

```



```

if resp[0] != 0x00:
    raise MifareError(
        f'Zapis Sector Trailer nieudany (status={resp[0]:#04x})')

```

Metoda `write_trailer()` rozmyślnie omija zabezpieczenie `write_block()` przed zapisem do trailerów. Zamiast tego bezpośrednio używa `InDataExchange` z komendą `WRITE` – po wcześniejszej walidacji bitów i autentykacji aktualnym kluczem. To celowe rozdzielenie: `write_block()` jest bezpieczny przez design, `write_trailer()` jest

świadomie wywołowaną operacją o potencjalnie nieodwracalnych skutkach.

Listing 3: main4.py – demonstracja zmiany kluczy

```

# main4.py – rozdział 4: zmiana kluczy i demonstracja zabezpieczeń

from machine import SPI, Pin
from pn532 import PN532
from iso14443 import ISO14443
from mifare_classic import MifareClassic, DEFAULT_KEY, KEY_A, KEY_B
from access_bits import (ACCESS_TRANSPORT, ACCESS_READ_A_WRITE_B,
                        validate_access_bits, describe_access)

spi = SPI(0, baudrate=1_000_000, polarity=0, phase=0,
          sck=Pin(18), mosi=Pin(19), miso=Pin(16))
czytnik_hw = PN532(spi=spi, cs_pin=17)
iso = ISO14443(czytnik_hw)
mifare = MifareClassic(iso)

# Nowe klucze – zmień na swoje własne przed użyciem!
NOWY_KLUCZ_A = bytes([0xA0, 0xA1, 0xA2, 0xA3, 0xA4, 0xA5])
NOWY_KLUCZ_B = bytes([0xB0, 0xB1, 0xB2, 0xB3, 0xB4, 0xB5])
SEKTOR = 2 # sektor testowy (nie używaj sektora 0!)

print('Rozdział 4 – zmiana kluczy MIFARE Classic')
print('UWAGA: Używaj tylko kart testowych z Zestawu!')
print('Przyłóż kartę...')
print()

karta = iso.wait_for_card(timeout_ms=15_000)
if karta is None:
    print('Brak karty.')
    raise SystemExit

print(f'Karta: UID={karta.uid_hex()} typ={karta.type}')

try:
    # 1. Waliduj bity dostępu przed zapisem
    ac = ACCESS_READ_A_WRITE_B
    print(f'Bity dostępu: {ac.hex(" ")} – walidacja: ',
          'OK' if validate_access_bits(ac) else 'BLAD!')
    for opis in describe_access(ac):
        print(f' {opis}')
    print()

    # 2. Zapisz nowy Sector Trailer w sektorze testowym
    print(f'Zmiana kluczy w sektorze {SEKTOR}...')
    mifare.write_trailer(
        sector_num=SEKTOR,
        key_a=NOWY_KLUCZ_A,
        access_bits=ac,
        key_b=NOWY_KLUCZ_B,
        auth_key=DEFAULT_KEY, # aktualny klucz (fabryczny)
        auth_key_type=KEY_A,
        uid=karta.uid,
    )
    print(' Klucze zmienione – zapisano Sector Trailer.')

    # 3. Weryfikacja: odczyt z nowym kluczem A (tylko odczyt)
    bloki = mifare.read_sector(
        SEKTOR, key=NOWY_KLUCZ_A, key_type=KEY_A, uid=karta.uid)
    print(' Weryfikacja odczytu nowym kluczem A: OK')

    # 4. Zapis z kluczem B
    dane = b'secured sector '
    mifare.authenticate(
        SEKTOR * 4, key=NOWY_KLUCZ_B, key_type=KEY_B, uid=karta.uid)

```

```
mifare.write_block(SEKTOR * 4, dane)
print(f' Zapis kluczem B: OK ({dane})')

except Exception as e:
    print(f'BLAD: {e}')
```

UWAGA: Używaj tylko kart testowych z Zestawu!

Program zapisuje nowe klucze do karty. Operacja jest nieodwracalna bez znajomości nowego klucza. Jeśli zmienisz klucze i zapomnisz ich wartości, sektor będzie niedostępny.

Do ćwiczeń z tym rozdziałem zawsze używaj jednej z kart testowych dołączonych do Zestawu do kursu RFID ze sklep.avt.pl – nie karty, której używasz w systemie produkcyjnym.

Jeśli przypadkowo zablokujesz sektor na karcie Gen1a: użyj backdoora 0x40 i wpisz znany Sector Trailer z powrotem.

Na oryginalnej karcie NXP – nie ma możliwości cofnięcia.

Wynik działania main4.py

Oczekiwany wynik w konsoli

Rozdział 4 – zmiana kluczy MIFARE Classic

UWAGA: Używaj tylko kart testowych z Zestawu!

Przyłóż kartę...

Karta: UID=04:A3:1B:7C:22:3F:80 typ=MIFARE Classic 1k

Bity dostępu: ff 0f 80 – walidacja: OK

Blok 0: R kluczem A lub B, W tylko kluczem B

Blok 1: R kluczem A lub B, W tylko kluczem B

Blok 2: R kluczem A lub B, W tylko kluczem B

Zmiana kluczy w sektorze 2...

Klucze zmienione – zapisano Sector Trailer.

Weryfikacja odczytu nowym kluczem A: OK

Zapis kluczem B: OK (b'secured sector ')

Materiały do dalszej lektury

Garcia et al., 'Dismantling MIFARE Classic' – oryginalna praca z 2008 roku, wciąż czytelna i konkretna. PDF dostępny na cs.ru.nl/~flaviog/publications/.

de Koning Gans et al., 'A Practical Attack on the MIFARE Classic' – szczegółowy opis Nested Attack z pomiarami czasu i skuteczności.

libnfc.org – biblioteka open-source do komunikacji z czytnikami NFC/RFID pod Linuxem; narzędzia mfoc i mfcuk działają na jej bazie.

Proxmark3 Wiki (github.com/RfidResearchGroup/proxmark3) – dokumentacja najpopularniejszego otwartego narzędzia do badań RFID.

NXP AN10922 – oficjalny dokument aplikacyjny opisujący protokół autentykacji MIFARE Classic.

Podsumowanie

Crypto1 to złamany algorytm – wiemy o tym od 2008 roku. Ale znajomość jego słabości to nie wyrok na karty MIFARE Classic w każdym zastosowaniu: systemy z prawidłowo skonfigurowanymi kluczami, bitami dostępu i dywersyfikacją kluczy per UID są wystarczająco bezpieczne dla wielu scenariuszy – biletów komunikacji miejskiej, kontroli dostępu do biur czy kart lojalnościowych. Problem dotyczy głównie systemów, które wyszły z linii produkcyjnej z kluczem fabrycznym.

Nasz stos kodu ma teraz cztery pliki po stronie Pico: pn532.py, iso14443.py, mifare_classic.py i nowy access_bits.py. W kolejnych rozdziałach rozszerzymy go o obsługę Ultralight (rozdział 5), NTAG (rozdziały 6...7) i NDEF (rozdział 8). Każdy z tych formatów jest prostszy niż MIFARE Classic pod względem bezpieczeństwa, ale ma własne cechy, które warto znać.

Pracownia konstrukcyjna

Kurs RFID składa się z 10 rozdziałów. Rozdział 1, prezentujący bazę sprzętową kursu, opublikowaliśmy w EP06/2026. Kolejne w bieżącym i kolejnym wydaniu. Listingi zostały napisane przez AI. Istnieje zgodna na świecie opinia, że zadania kodowania najlepiej wykonuje Claude. Robi to (podobno) bezbłędnie, w związku z tym ogłoszono koniec profesji junior programista. Sprawdźmy to na sobie. Ten kurs jest wersją nie zweryfikowaną. Potrzebny do tego kursu sprzęt skompletowaliśmy jako „Zestaw do kursu RFID”, wkrótce dostępny w sklep.avt.pl. Pięć takich zestawów prześlemy bezpłatnie Czytelnikom, którzy podejmą się weryfikacji listingu i przekazania uwag do redakcji EP (redakcja@ep.com.pl). Czekamy na zgłoszenia, w których prosimy przedstawić dotychczasowe doświadczenie w elektronice embedded.



Jak (dobrze) projektować elektronikę użytkową?

Do napisania tego tekstu skłoniły mnie problemy z pewnym urządzeniem, z którego jedno z moich dzieci korzysta codziennie w szkole, a które kosztuje 18 tysięcy złotych. Problemy te są skutkiem idiotycznych decyzji inżynierów, którzy owo urządzenie zaprojektowali, stawiając na obniżenie kosztów produkcji pomimo niebagatelnej marży, którą producent sobie nalicza. A wystarczyło zmienić kilka drobiazgów i wziąć pod uwagę docelowe warunki pracy, by uniknąć tych problemów.

Owym tajemniczym urządzeniem jest powiększalnik. Ale nie klasyk fotografii analogowej, lecz przyrząd ułatwiający niedowidzącemu dziecku funkcjonowanie w szkole. Model Exigo, o którym będzie mowa (**fotografia 1**), posiada dwa ekrany, jeden do powiększania tego, co znajduje się na blacie, czyli zazwyczaj podręcznika lub zeszytu, drugi ekran zaś pokazuje to, co znajduje się na szkolnej tablicy lub co robi nauczyciel. Niemal wszystkie powiększalniki na rynku mają tylko jeden ekran i jeśli posiadają dodatkową kamerę do dali, to obrazy z obu kamer są wyświetlane na jednym ekranie. Pod tym względem Exigo jest urządzeniem unikalnym, bo znam tylko dwa inne modele posiadające dodatkowy ekran, jednym z nich jest CloverBook Pro, który posiada jeden ekran, ale drugi można dokupić i przyczepić do niego. Pozwala też na podłączenie dowolnego innego ekranu dzięki gniazdu HDMI. Exigo jednak ma dwa całkowicie niezależne ekrany i jest to jedyne takie urządzenie na świecie. I są ku temu dobre powody.

Błędy projektu powiększalnika Exigo, czyli jak zrobić urządzenie na (dwa) lata?

Powiększalnik Exigo jest wykonany dość solidnie. Po rozłożeniu nóg stawia się go na blacie i dolny ekran można przesuwac na boki, co może pomóc przy czytaniu i pisaniu w dużym powiększeniu. Drugi ekran jest umocowany na dużym przegubie kulowym. I tu są pierwsze dwie wady: dolny ekran lata na boki luźno, chyba że zakręci się pokrętło „hamulca”. W pozycji transportowej pokrętło to ukryte jest za przegubem kulowym. Producent zatroszczył się o to, by mechanizm wózka dolnego ekranu nie uderzał brutalnie o końce statywu, wklejając tam gąbki. Gąbki te zniknęły po paru miesiącach normalnego użytkowania w klasie. Sam przegub też wymaga zaciskania, bo inaczej górny ekran będzie się majtał na wszystkie strony. Problem w tym, iż nie zawsze korzysta się z górnego ekranu, więc

szanse na to, iż dziecko go zablokuje, są znikome. Stopy, na których Exigo stoi, są wykonane z twardej gumy, ale pochodzą z zestawu standardowych elementów gumowych, przez co ich zaokrąglone spody dotykają blatu niewielką tylko powierzchnią. Stopy w moim statywie fotograficznym za 80 złotych są wykonane na kulowych przegubach, przez co zawsze dotykają podłoża całą swoją płaską powierzchnią, nawet jak podłoże jest nierówne. Szkoła, wiedząc, jak wygląda życie klasowe w podstawówce, ubezpieczyła urządzenie Exigo na wypadek nieszczęśliwych zdarzeń, co jest znakomitą pomysł, biorąc pod uwagę te podstawowe niedoróbki mechaniczne. Same ekrany zabudowane są dość solidnie w korpusy z tworzywa i gumy o grubości przynajmniej 4 mm. Ekranami tymi są tablety Samsung Galaxy Tab S, na których zainstalowane jest oprogramowanie Exigo. Dolny korpus posiada własny układ optyczny z kamerą umieszczoną centralnie i diodami doświetlającymi blat. Taśma MIPI biegnie do narożnika tabletu, gdzie wpięta jest w miejsce oryginalnej kamery. Skąd to wiem? Bo w tym miejscu korpusu jest zaślepka umocowana „na wcisk”, którą znudzony życiem użytkownik może wyciągnąć paznokciem. Górny ekran w miejscu zaślepki ma własny obiektyw do dali. To akurat sensowne rozwiązanie, bo standardowe obiektywy kamer w tabletach i smartfonach mają małe rozmiary, a przez to równie małą jasność – w gorszych warunkach oświetleniowych polskiej klasy standardowy obiektyw byłby za ciemny. Wspomniane podświetlenie zasilane jest z akumulatora ukrytego w wózku. Między wózkiem a modulem z kamerą biegnie przewód, bo cały ekran odchyła się względem wózka. Inżynier nie przewidział, iż ten przewód w codzienności szkolnej podlega podwyższonemu ryzyku wyrwania. Włącznik podświetlenia znajduje się obok pokrętła hamulca wózka. Gniazdo ładowania to klasyczne gniazdo DC znane ze starszych laptopów. Gdzie je umieszczono? Z tyłu wózka, tuż obok przegubu.



Fotografia 1. Powiększacz dwuekranowy Exigo za ponad 18 tysięcy złotych (a) oraz przykład użycia w klasie (b)

Kabel ładowania powinien mieć wtyczkę kątową, by nie uderzać o przegub i nie uszkodzić siebie ani gniazdka. Czy jest kątowy? Oficjalnie nie. Inżynier nie pomyślał o tym, iż przy luźnym hamulcu ekran z wózką mogą się przemieszczać wte i nazad, a kabel ładowania podświetlenia będzie uderzał o przegub, wydatnie skracając jego żywotność. W efekcie po niespełna dwóch latach podświetlenie się już nie ładuje. Nie dość bowiem, iż wewnątrz kabla nastąpiło zwarcie (które uszkodziło ładowarkę), to jeszcze samo gniazdko ładowania uległo skutkom mechanicznych naprężeń i nie „styka”. Jak wspomniałem, „sercami” Exigo są dwa tablety. By nikt nie wykradł sekretów oprogramowania (jakby było co kraść), gniazda USB-C do ładowania tabletów są niedostępne. Zamiast tego są bajeranckie złączki magnetyczne oraz równie bajeranckie przewody ładowania z końcami świeącymi na niebiesko – najbardziej irytujący i nadużywany kolor diod LED w elektronice użytkowej. Praktycznie od nowości były problemy z tymi złączkami – odrobina jakiegokolwiek brudu spowalnia ładowanie. Obecnie, z powodu uszkodzenia ładowarki tablety te ładują się z prędkością <10% na godzinę. A rozładowują się o 15...20% na godzinę lekcyjną. Dlatego konieczny będzie zakup nowej ładowarki, ale też doładowywanie urządzenia z zewnętrznego powerbanku w trakcie zajęć lekcyjnych. Dodajmy do tego kolejny fakt: typowa bateria w tablecie wytrzyma od dwóch do pięciu

lat takiego traktowania. Potem puchnie. Zatem, podsumowując, mamy dwa tablety za mniej niż pięć tysięcy złotych uwięzione w kiepsko zaprojektowanej konstrukcji za kolejne trzynaście tysięcy, przy czym najdalej po pięciu latach oba tablety będą się nadawać do elektrośmieci. Porównajmy to z konstrukcją CloverBook Pro (fotografia 2). Mamy tu większy ekran dotykowy (12,4”), fizyczne przyciski i pokrętła do kontroli urządzenia, rozkładany statyw niewymagający pokręteł i hamulców, drugi ekran dopinany do pierwszego standardowym przewodem HDMI, zintegrowane podświetlenie blatu i kamerę do dali oraz wymienną baterię. Do ładowania służy standardowe gniazdo DC i nie jest ono narażone na naprężenia boczne. Ba, nawet torba transportowa jest mniejsza i lepiej zaprojektowana. Jedyna wada? Niedostępne w Polsce. Model z dodatkowym ekranem kosztuje 3344 USD, czyli w chwili pisania tych słów 12 227 zł. Jest jeszcze większy model, CloverBook Pro XL, z dwoma ekranami 16” za 4190 USD, czyli 15 320 złotych, wciąż mniej niż kosztuje Exigo ze swoimi dwoma tabletami 10,5”. Rozważam sprowadzenie któregoś z tych urządzeń do Polski na własny koszt, bo nie wierzę, by Exigo dotrwało do końca podstawówki przy tak nieprzemysłanej konstrukcji.

Jak nie robić elektrycznego piekarnika dla niewidomych i niedowidzących?

Amerykański start-up, Tovala, wypuścił jakiś czas temu stołowy piekarnik elektryczny zaprojektowany z myślą o osobach niewidomych i niedowidzących. A raczej sprzedaje nie tyle piekarnik, co ekosystem oparty na tym piekarniku. To sprawia, iż sam piekarnik jest lekko „niedorobiony”. Tovala Smart Oven to kolejny produkt „smart”, którego funkcjonalność opiera się na aplikacji i chmurze. Bez aplikacji sam piekarnik jest prawie bezużyteczny.



Fotografia 2. CloverBook Pro, przykład jak dobrze zaprojektować powiększacz do szkoły

Z aplikacją użyteczność rośnie, ale tylko i wyłącznie dlatego, iż system operacyjny w smartfonie (Android lub iOS) ma funkcje ułatwień dostępu. Bez aplikacji osoba niewidoma nie jest w stanie ustawić parametrów piekarnika. Tylko wybór trybów pracy jest możliwy z panelu przedniego, bo każdy tryb ma swój przycisk. Temperatura i czas pracy też są ustawiane przyciskami, ale użytkownik nie dostaje żadnej informacji zwrotnej, bo Towała wbudowała w piekarnik Wi-Fi, skaner kodów i różne funkcje „smart”, ale nie wbudowała najbardziej nawet prymitywnego syntezatora mowy (sic!). Serio, Towała?! O układzie TMS5100 od TI nie słyszeliście? Moje chińskie słuchawki Bluetooth za 40 złotych przemawiały do mnie ludzkim głosem. Dlaczego firma Towała zapomniała o tym podstawowym ułatwieniu dostępu dla niewidomych? Cóż, zajęci byli rozmyślaniami, jak z tych samych niewidomych wyciągnąć więcej pieniędzy na usłudze dostarczania gotowych posiłków. Rozwiązanie dość proste: klient dostaje paczkę z jedzeniem, skanuje kod QR aplikacją w telefonie albo za pomocą samego piekarnika, a wszystkie parametry same się ustawiają. Aplikacja wyświetli (i odczyta za pomocą systemowego silnika TTS) instrukcję przygotowania, opis produktu i wartości odżywcze. Nie wiem jednak, czy płacenie odpowiednika 38..50 zł za posiłek dla jednej osoby ma sens. Co prawda, jest to tańsze niż w amerykańskiej restauracji, ale i tak droższe niż samodzielne gotowanie. Dodatkowo Towała nie dopracowała nawet tych posiłków, bo czasami trzeba zeskanować więcej niż jeden kod QR, albo są trzy kody do wyboru dla jednego posiłku, i trzeba wybrać jeden z nich zależnie od głównego składnika. Oznaczenia na opakowaniach sosów czy innych dodatków są nadrukowane małą czcionką – aplikacja w rodzaju Seeing AI sobie poradzi, ale niewidzący już nie. Z ciekawości pobrałem też instrukcję obsługi w formacie PDF – wygląda jak

typowa instrukcja dla drobnego AGD. Nie jak instrukcja dla sprzętu stworzonego z myślą o osobach niedowidzących. Mój brat jest niewidomy. Lubi też gotować. Jego kuchenka ma mechaniczne pokrętki trybów grzania i mechaniczne pokrętki temperatury. Timer elektroniczny, ale on to obchodzi, używając asystenta Google albo Siri, by ustawić sobie minutnik. W Wielkiej Brytanii z kolei jest firma produkująca udźwiękowiony sprzęt AGD. Urządzenia w stylu retro z mechanicznymi pokrętkami też są popularne, bo każdy może zapamiętać pozycję gałki dla zadanej temperatury czy zadanego trybu. Na rynku od lat dostępne są mówiące wagi kuchenne i łazienkowe, mówiące termometry czy mówiące naczynia miarowe. Te ostatnie są dość drogie, bo to bardzo niszowy produkt. Technologia „udźwiękowienia” sprzętu nie jest skomplikowana, gdyż były produkowane dedykowane układy scalone syntezytorów, jak wspomniany TMS5100 z lat 80. Artykuły na ten temat ukazały się na łamach „Elektroniki Praktycznej” w numerach 3/97, 4/97 i 5/97. Zestaw nosił oznaczenie AVT-322 i opierał się na wykorzystaniu scalonego rejestratora mowy ISD2560 (co oznaczało konieczność nagrania odpowiednich słów samodzielnie). W podobnym czasie miałem też elektroniczny zegarek naręczny z syntezą mowy, którego największą wadą był, moim zdaniem, odgłos budzika – niezbyt udane „Kukuryku!”. Interfejs (nadmiernie) dotykowy i inne problemy

Nigdy nie zrozumieję trendu zamieniania fizycznych przycisków na pola dotykowe. Upraszcza to w teorii produkcję obudowy, ale w praktyce to nie ma znaczenia, bo i tak trzeba zrobić formę wtryskową, a otwory i wycięcia w takiej sytuacji nie generują dodatkowych kosztów. Koszty generuje na przykład plastikowa wkładka z klawiszami oraz same fizyczne przyciski, i dla oszczędzenia tych paru groszy (co nie zmienia ceny dla klienta, ale zmniejsza o ułamek procenta marżę producenta) firmy rezygnują z rozwiązania dobrego na rzecz rozwiązania irytującego. Przykładem może być osuszacz powietrza, który stoi obok biurka, przy którym piszę te słowa. Osuszacz całkiem dobry, ale jego interfejs opiera się na pięciu przyciskach, których ja nie widzę, bo są dotykowe i dość małe. Gdyby miały podświetlenie, byłoby lepiej, no ale te przyciski go nie mają. Producent nie pomyślał nawet o tym, by umieścić wokół nich lekko wypukłe ramki ułatwiające ich zlokalizowanie „po omacku”. Jeśli więc chcę wyłączyć osuszacz, to głaszczę go ręką przy jednej krawędzi w nadziei trafienia na pierwszy przycisk dotykowy. No ale wyłączanie i włączanie osuszacza to tylko połowa problemu. Drugą połową są moje koty, które używają go jako drogi na blat biurka. Za każdym razem któreś ze zwierząt dotyka jednego z pól i albo włącza osuszacz, gdy miał być wyłączony, albo zmienia próg aktywacji lub inne parametry. Każdy, kto ma kota, wie, iż łatwiej jest przybić galaretkę do ściany, niż wytresować kota, by czegoś nie robił lub gdzieś



nie łąził. Skazany więc jestem na monitorowanie osuszacza, aż do śmierci mojej, kotów lub osuszacza właśnie. Moje inne słuchawki Bluetooth miały pola dotykowe do kontrolowania odtwarzania i odbierania połączeń. Ze zdziwieniem odkryłem, iż jeśli się spociłem z wysiłku, to moje wilgotne włosy zatrzymywały mi muzykę albo kończyły połączenie, jeśli przypadkiem dotknęły któregoś z tych pól. Innym problemem współczesnego sprzętu domowego jest zubażanie interfejsów na rzecz aplikacji dla smartfona. Nie lubię instalować aplikacji tego typu. Nie chcę kuchenki, pralki czy robota kuchennego sterowanego telefonem. Nie chcę odkurzacza, którego nie da się włączyć fizycznym przełącznikiem. Nie jestem zainteresowany instalacją dodatkowych aplikacji byle tylko móc używać sprzętu codziennego użytku. Nie instaluję też aplikacji bankowych, do szybkich przelewów, do tańszych zakupów w różnych sieciach sklepów czy nawet aplikacji mObywatel. Moje życie codzienne nie kręci się wokół smartfona, którego najbardziej zaawansowaną funkcją używaną przeze mnie codziennie jest zdolność odtwarzania audiobooków (i to nie z jakiejś aplikacji opartej na abonamentach czy inne głupoty – pliki audio mam na karcie microSD). Kilka lat temu kupiłem lampki choinkowe z funkcją IoT. Zainstalowałem aplikację Tuya i przez jeden okres świąteczny fajnie było wybrać kolor i tryb świecenia z aplikacji. Ale potem zmienił mi się router i w następnym roku już to nie działało, a grzebać w ustawieniach aplikacji mi się nie chce. Za to Tuya nieustannie przesyła mi wiadomości e-mail o zmianach w regulaminie i inne głupoty związane z aplikacją, której od dwóch lat nie włączyłem. IoT to rozwiązanie desperacko poszukujące problemu do rozwiązania. Nie mam prawa jazdy z dość oczywistych względów, ale co nieco wiem o motoryzacji. Dlatego doznałem znacznego zadziwienia, gdy dowiedziałem się jakiś czas temu, iż producenci samochodów, głównie w USA, rezygnują z tradycyjnego hamulca ręcznego na rzecz „elektronicznego hamulca ręcznego”. Dość szybko pojawiły się doniesienia o jego problemach – w razie awarii elektryki samochodu hamulec ten nie działa wcale. Kilka osób miało problem z zatrzymaniem pojazdu na pochyłej drodze, a inne nie mogły odblokować hamulca w celu odholowania do serwisu. Nie wiem, jak czytelnicy, ale moim zdaniem w każdym samochodzie dźwignia hamulca ręcznego i kierownica powinny zachować mechaniczne połączenie z hamulcami i mechanizmem kierowniczym, właśnie na wypadek awarii. Dlatego więc producenci zamiast odłączać całkowicie hamulec ręczny, zwyczajnie nie dodadzą do niego wspomagania tak, jak dodają wspomaganie kierownicy czy sterownik ABS? Że drogo? Spokojnie, klient zapłaci. Amerykańskie samochody są generalnie dziwne – wszystko jest podpisane (po angielsku), podczas gdy reszta świata już wiele dekad temu przerzuciła się na ikonki. Kolejna, głupia rzecz we współczesnej motoryzacji to zastępowanie wszystkich

fizycznych przełączników i regulatorów tabletem wkomponowanym w deskę rozdzielczą. Takie rozwiązanie sprawia, iż kierowca musi odrywać wzrok od drogi za każdym razem, gdy chce zmienić ustawienia klimatyzacji czy zmienić stację radiową czy odtwarzaną muzykę. W starszych modelach z mechanicznymi kontrolkami klimatyzacji i tradycyjnym radiem samochodowym wszystkie te czynności można robić po omacku. Nowoczesne rozwiązanie zatem staje się czynnikiem ryzyka wypadku. Z tego samego powodu w różnych krajach obowiązują zakazy używania smartfona w trakcie prowadzenia samochodu, a niejeden wypadek był spowodowany stukaniem w ekran takiego smartfona. Pomijam już fakt, iż sprzęt w tych „tabletach” nie jest najwyższej klasy, zwykle w modelach budżetowych. Osobne pytanie: jak dobrze są zabezpieczone te tablety? Czy ktoś z zewnątrz nie mógłby uzyskać zdalnego dostępu do tabletu, a stamtąd przez wewnętrzną sieć samochodu na przykład dobrać się do ECU? Biorąc pod uwagę, iż znaleziono luki bezpieczeństwa w rozrusznikach serca i pompach insuliny, to się zastanawiam nad jakością oprogramowania w dużo bardziej złożonych systemach samochodów.

Trwałość i potencjał do naprawy

Przez lata producenci oraz zatrudnieni przez nich lobbyści próbowali zrobić wszystko, by z jednej strony odebrać ludziom prawo do naprawy, a z drugiej wmówić im, iż nie ma czegoś takiego jak planowe psucie. Oczywiście nie trzeba tworzyć jakichś specyficznych pułapek w oprogramowaniu – wystarczy tak dobrać komponenty lub materiały, by po określonym czasie nastąpiła awaria, najlepiej dwa dni po zakończeniu okresu gwarancyjnego. Nie zmienia to faktu, iż na przykład Apple zostało przyłapane na tworzeniu aktualizacji iOS, które gorzej działają na starszych wersjach smartfonów iPhone. Jednocześnie same smartfony, i to nie tylko produkty Apple, stają się nienaprawialne. Ba, wymiana akumulatora litowego wymaga dosłownego rozklejenia urządzenia. Mam tablet, w którym muszę wymienić ekran. Tylny korpus dało się zwyczajnie od urządzenia odpiąć, ale sam panel jest przyklejony do ramki z elektroniką za pomocą wyjątkowo mocnej taśmy dwustronnej. Również akumulator jest przyklejony do ekranu, i nieco się obawiam uszkodzenia jego hermetycznego opakowania. Ostatnio też opaska sportowa mojej żony po paru miesiącach od wymiany ogniów nagle i spontanicznie odmówiła posługi. Po krótkiej diagnozie problem wynikł po części z mojej winy – na uszczelce, przypuszczam, że znalazła się odrobina brudu, co spowodowało, iż opaska przestała być w stu procentach szczelna. Niedoróbką producenta jednak była oszczędność na materiałach, gdyż przy skręcaniu obudowy (po zrobieniu zdjęć na prośbę dystrybutora) plastikowy element obudowy, w który wchodzi wkręt, zwyczajnie pękł. Wcześniejszy model stosował wtapiane tuleje

gwintowane, co znacząco podnosiło trwałość, oraz dodatkową, wewnętrzną osłonę elektroniki. Ten model został wycofany, a nowszy uległ procesowi „racjonalizacji”, czyli znanemu z czasów PRL-u procesowi cięcia kosztów i oczekiwani. Baterie w obu modelach zmienia się po około roku, zakładając nawet najostrożniejsze obchodzenie się z obudową i uszczelką przypuszczam, że po trzech lub czterech wymianach ten sam element obudowy pęknie. I nie mówimy tu o jakimś chińskim produkcie, lecz o opasce sportowej bardzo znanej i uznanej firmy. Marki i modeli nie podaję, bo jestem po weekendowym kursie unikania pozwów. Nie mogę udowodnić celowości ich działań, więc jakbym podał więcej danych identyfikacyjnych, to ja bym im płacił karę zarządzoną sądownie, a nie oni mi (i innym swoim klientom). Ciekawy i na moje oko raczej nienaprawialny problem miałem z dość drogim skanerem ręcznym. Nie wiem, co się z nim konkretnie stało, bo był odpowiednio spakowany, ale po przeprowadzce przestał poprawnie skanować, zostawiając czarny pas szeroki na 1/5 strony. Parę miesięcy temu postanowiłem go otworzyć i spróbować naprawić. Teraz 3/5 szerokości to czerń. Wydałem więc 45 złotych na znanym portalu aukcyjnym na używany wizualizer, który wykorzystuje kamerę do skanowania. Co ciekawe, lekko poobijane urządzenie sprzed kilku lat, z brakującym kawałkiem obudowy wciąż działa i robi to, czego można od niego oczekiwać. Dlaczego? Cóż, ręczny skaner przeznaczony był dla studentów i użytkowników domowych. Wizualizer zaś produkowany był z myślą o szkołach, uczelniach oraz środowisku biznesowym. Stąd solidna, gruba obudowa i spora masa własna z wieloma elementami metalowymi. Biorąc pod uwagę kształt i masę, mógłbym użyć wizualizera w obronie własnej, a potem przed sądem zrobić prezentację, dlaczego to była obrona konieczna. To jest właśnie ciekawa różnica między sprzętem dla zwykłego konsumenta, a sprzętem dla biznesu i rynku profesjonalnego. Mam biurową drukarkę laserową, która działa relatywnie sprawnie już ponad dwie dekady. Wymagałaby zapewne czyszczenia i konserwacji, ale na tym kompletnie się nie znam. W tym czasie miałem też trzy kolorowe „plujki”, przy czym w ostatniej mimo ustawienia „do transportu” głowica się zatkała – producent napisał w instrukcji, iż trzeba na niej drukować przynajmniej raz na tydzień wszystkimi kolorami. Moja pierwsza „plujka” nie stawiała takich wymagań, w zamian trzeba było kupować kartridże z głowicami w cenie połowy nowej drukarki. Dlaczego w modelu 15 lat starszym głowice były trwalsze i się nie zatykały, a w nowszym już był problem? Przypuszczam, iż winne jest cięcie kosztów. Myślę, iż następnym razem znów kupię urządzenie przeznaczone do pracy w biurze, a nie w domu. Zapłacę więcej, ale znów będę miał spokój na 20 lat. Wracając na moment do sprzętu AGD, to nauczony doświadczeniem przestałem kupować najtańsze urządzenia. Mój primer

blender, marki Kenwood, okazał się nietrwały i uległ awarii z powodu pęknięcia plastikowego elementu wewnątrz. Dołączony do zestawu młynek do kawy był tak źle zrobiony, iż ostrza starły wewnątrz pojemnika z tworzywa. A wystarczyło lepiej dopasować wymiary i użyć nieco grubszego tworzywa w mechanizmie zabezpieczającym. Blender od dość znanego polskiego producenta, marki Yoer, okazał się być dużo lepiej zaprojektowany, choć kosztował tyle, co półtora Kenwooda. Przy okazji przetestowałem dostępność części zamiennych po tym, jak jeden z moich kotów strącił pokrywkę od młynka. Nową pokrywkę udało się zakupić bez problemu. Swoją drogą firma Kenwood już na zawsze straciła jednego potencjalnego klienta, gdyż ich budżetowy produkt okazał się zbyt budżetowy i nietrwały – nie zaufam im nawet jeśli mają w ofercie trwałe, acz nieco droższe sprzęty.

Dobry inżynier, zły księgowy, chytry prezes...

Racjonalizacja kosztów, atrakcyjny projekt i chęć unowocześniania wszystkiego na siłę są prostą drogą to tworzenia produktów, które są awaryjne, a czasami wręcz niebezpieczne. Dobry inżynier zakłada, w jakich warunkach jego urządzenie będzie używane, i przez kogo. Lepszy inżynier organizuje beta testy na klientach i nie w formie sprzedaży gotowego produktu i analizy raportów z serwisu, lecz zapraszając tych klientów do swojej pracowni. W ten sposób Tovala by nie wypuściła piekarnika dla niewidomych, który nie mówi, a w powiększalniku Exigo ładowanie podświetlenia byłoby umieszczone gdzieś indziej, lub wtyk ładowarki zostałby zamieniony na wariant kątowy. Dobry inżynier nie zamieniłby elementu krytycznego dla bezpieczeństwa użytkownika z rozwiązania stricte mechanicznego na rzecz rozwiązania stricte elektronicznego. Dobry inżynier, jeśli musi zaprojektować urządzenie, które zepsuje się po gwarancji, projektuje margines większy niż dwa dni po gwarancji, bo rozrzut produkcyjny może być duży. I robi to tak, by nie było zbyt oczywiste dla innych. Lepszy inżynier projektuje urządzenia, które są zwyczajnie trwałe. Dobry inżynier zakłada najgorszy możliwy scenariusz, i projektuje adekwatne rozwiązania na scenariusz jeszcze dwa razy gorszy. Dobry inżynier projektuje interfejs użytkownika, który będzie dobry dla każdego użytkownika, także niepełnosprawnego. I nie zastępuje go apką na smartfona. Dobry inżynier wie, iż warto otaczać pola dotykowe wypukłą otoczką, zwłaszcza że w produkcji to nic nie kosztuje. Producenci sprzętu powinni się już nauczyć, iż łatwo jest zrazić do siebie klienta, i bardzo trudno odzyskać jego zaufanie. Konkurencja na rynku jest duża, więc klient zawsze może wybrać innego producenta. Tu przychodzi mi na myśl mój epizod z kartą graficzną ATI, która z powodu niedoróbki po stronie sterownika (problem z błędami mikrokodu kontrolującego pracę

karty, a wczytywanego przez sterownik), karta ta wieszła się pod obciążeniem. Pozbyłem się jej i już nigdy nie kupię żadnego Radeona i nie zamierzam kupować. Generalnie nie jestem wybrednym klientem, oczekuję minimum trwałości i wytrzymałości oraz bezproblemowej pracy. Jeśli produkt nie spełni tych oczekiwań, to już więcej nie kupię nic od tego producenta. Nie zawsze jednak dobry inżynier ma wybór, bo nad nim stoi księgowy-liczykrupa oraz chytry, bezduszny prezes. Problemem dzisiejszego świata jest tendencja wielkich korporacji do narzucania rynkowi produktów, które są niskiej jakości. Przykład idzie z USA, gdzie przez ostatnie dekady wszyscy obywatele kupowali wszystko na kredyt. Amerykanie są namawiani do wyrobienia karty kredytowej już w liceum, a punktacja oceny zdolności kredytowej jest niemal miarą zdrowej postawy obywatelskiej i patriotycznej. Im więcej punktów, tym korzystniejsze warunki kredytowe i tym większa kwota kredytu, ale nie tylko. Zdolność kredytowa określa możliwości wynajęcia lokalu mieszkalnego, jakość dostępnego ubezpieczenia zdrowotnego czy nawet dostęp do wyższej edukacji. Dlatego Amerykanie kupują wszystko na kredyt, a producenci sprzętu celowo obniżają jakość i trwałość by napędzać spiralę konsumpcjonizmu. Dlatego Apple stworzyło swego czasu aktualizację systemu, która celowo gorzej działała na starszych modelach smartfonów i tabletów (za co zostali ukarani). Lobbyści amerykańskich korporacji aktywnie walczą z prawem do naprawy nie tylko w USA, ale też w UE (szerzej o prawie do naprawy pisałem w EP 05/2026). Tymczasem Amerykanie, zwłaszcza ci zarabiający w okolicy średniej krajowej lub mniej, dość łatwo wpadają w spiralę kredytową, z której jest im bardzo ciężko wyjść. Średnie zadłużenie per capita w USA wynosi ponad 63 tysiące dolarów, czyli niemal 235 tysięcy złotych. Dla porównania średnie zadłużenie per capita w Polsce wynosi maksymalnie 25 tysięcy złotych. Na tę przepaść składają się głównie długi medyczne, ale też fakt, iż przeciętny dorosły Amerykanin ma od 3,7 do 4 kart kredytowych. Średnia stopa oprocentowania tych kart wynosi 20...21% w skali roku. W ostatnim czasie jednak w USA konsumpcjonizm hamuje, a obywatele mają coraz większe oczekiwania względem trwałości i jakości produktów. Amerykańska gospodarka hamuje, bezrobocie rośnie, koszty życia, a zwłaszcza energii też. UE nie poddaje się w kwestii prawa do naprawy. Trendy się zmieniają, i coraz więcej konsumentów myśli tak, jak ja. Jest zatem szansa, iż dobrzy inżynierowie pokonają złych księgowych i chytrych prezesów. Miejmy nadzieję, że tak się stanie.

Paweł Kowalczyk

Odnosiniki:

[1] – <https://www.youtube.com/watch?v=uBuowtmyLLO>



Miesięcznik „Elektronika Praktyczna” (10 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:

Wiesław Marciniak

Sekretarz Redakcji:

Dariusz Welik

Menedżer Magazynu:

Katarzyna Gugała, katarzyna.gugala@ep.com.pl

Szef Pracowni Konstrukcyjnej:

Jakub Sobański

Zespół marketingu i reklamy:

Katarzyna Gugała, Bożena Krzykawska, Grzegorz Krzykowski

Stali współpracownicy:

Paweł Kowalczyk, Michał Kurzela, Jakub Tyburski

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 AVT Korporacja 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”.



KURS PRAKTYCZNY AI

Praktyczne podejście. Zero marketingowej mgły!



Zamów na [UlubionyKiosk.pl](https://ulubionykiosk.pl)

