

ELEKTRONIKA PRAKTYCZNA

● Międzynarodowy magazyn elektroników konstruktorów ● czerwiec ● 6/2026 ●



AI w pracowni elektronika konstruktora

Z praktyki ekspertów

- Zephyr RTOS – pierwsze 48 godzin: od Blinky do BLE
- Półprzewodniki szerokoprzerwowe GaN vs. SiC – który wybrać do przetwornicy 200 W?
- Analiza, symulacja i eksperyment – klucz do udanego projektu kompensacji zasilacza

W pracowni

- Cyfrowy nadajnik FM
- Nagrzewnice indukcyjne ZVS
- Zegary hobbystyczne

Moduły z Chin

- M5Stack Cardputer ADV z modułem LoRa-1262

Repetitorium

- Oscyloskopy 2026
- Wzmacniacze operacyjne i komparatory

Kity AVT

- Elektroniczna klepsydra LED (AVT6098)

Podzespoły

- Przekazniki półprzewodnikowe SSR – budowa, działanie i zastosowanie

Audio bez tajemnic

- Syntezatory dźwięku. Filtry sterowane napięciem

AI w elektronice

- Sprzęt, który musiał dorosnąć

Kurs

- Kurs RFID z PN532, rozdział 1. RFID od zera: podłączamy PN532 i piszemy pierwszą klasę



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

AVT-Korporacja 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

Klaruje się EP po zmianach

Koszt druku ograniczał objętość EP. Po rezygnacji z wersji drukowanej objętość wzrosła dwukrotnie. Rozprostowaliśmy skrzydła do wyższych lotów. Przede wszystkim zwiększyliśmy ilość materiałów tutorialowych – repetytoriów, kursów i artykułów z cyklu „z praktyki ekspertów”, które prowadzą Czytelnika krok po kroku przez konkretny problem konstrukcyjny. Odpowiedzi na Ankietę upewniły nas w słuszności podjętych zmian. Ogólnie Czytelnicy ocenili zmienione EP na 4+. Dziękujemy za gremialny udział w Ankiecie, szczególne wyrazy wdzięczności kierujemy do tych Czytelników, którzy napisali do nas listy, niekiedy bardzo obszerne, z wieloma uwagami i sugestiami. Na następnej stronie publikujemy najbardziej charakterystyczny przykład takiej korespondencji. Feedback działa i klaruje profil tematyczny i formę EP. Nie traktujemy nowej formuły jako zamkniętej – wciąż ją dostrajamy, a Państwa głosy są w tym procesie najważniejszym instrumentem. W tym numerze dwie sprawy zasługują na szczególną uwagę.

1. Projekty. Wielu Czytelników domaga się projektów w EP. Przed trzydziestu laty atrakcyjne projekty zbudowały popularność EP i kitów AVT. Jednak świat się zmienił. Teraz pasjonaci majsterkowania (makers według obecnej terminologii angielskiej) mają tysiące opisów projektów w internecie, w tym kilka tysięcy z dorobku AVT. Samo dorzucenie do tej masy kolejnego schematu niewiele wnosi. Doszliśmy do wniosku, że rolą redakcji EP powinno być tworzenie syntetycznego przewodnika po określonych kategoriach tematycznych projektów – pokazanie, czym różnią się dostępne rozwiązania, jakie kryją pułapki i jak wybrać to właściwe dla własnych potrzeb. Proponujemy następującą formułę działu W PRACOWNI: artykuł składa się z dwóch części A i B. W części A prezentujemy przykładowy projekt, a część B zawiera krytyczne komentarze do tego projektu oraz systematyzuje kategorie projektów pokrewnych, dostępnych w internecie. W bieżącym numerze ten kierunek ilustrują artykuły o cyfrowym nadajniku FM, nagrzewnicach indukcyjnych ZVS oraz zegarach hobbystycznych. Czekamy na Państwa opinie, czy taka formuła trafia w oczekiwania. Zależy nam zwłaszcza na części B – to w krytycznej analizie i uporządkowaniu pokrewnych rozwiązań, a nie w samym schemacie, widzimy dziś największą wartość dla konstruktora.

2. AI w pracowni elektronika. To temat arcyważny. Od AI nie uciekniemy. To temat okładkowy tego numeru. Sytuacja rozwija się bardzo dynamicznie. Próbuje ocenić stan rzeczywisty według najświeższych publikacji z ostatnich miesięcy – bez hurraoptymizmu, ale i bez lekceważenia. Najbardziej gorący spór toczy się wokół przydatności AI do kodowania elektroniki embedded. Jedni przekonują, że modele językowe piszą już kod sprawniej niż przeciętny inżynier – programista; inni wskazują, że potrafią one wstawić lukę bezpieczeństwa albo „wyhalucynować” nieistniejący rejestr czy bibliotekę. Zamiast rozstrzygać ten spór zza biurka, proponujemy eksperyment. Sprawdzmy to na sobie. W Pracowni Konstrukcyjnej AVT opracowano Kurs RFID z układem PN532, w którym listingi są dziełem AI; jego pierwszy rozdział znajdują Państwo w bieżącym numerze. Zapraszamy Czytelników do weryfikacji tych listingów – do sprawdzenia, czy kod się kompiluje, czy działa na docelowym sprzęcie i czy nie kryje błędów, których model nie zauważył. Przygotowaliśmy 5 zestawów sprzętu do realizacji tego kursu. Rozdamy je bezpłatnie Czytelnikom, którzy podejmą się takiej weryfikacji, a jej wyniki – także te krytyczne – opublikujemy w jednym z kolejnych wydań. Proszę o listowne zgłoszenia pod adresem redakcja@ep.com.pl.

Poza tymi dwoma wątkami numer jest wyjątkowo obfity. W dziale „z praktyki ekspertów” piszemy m.in. o wyborze między półprzewodnikami GaN i SiC w przetwornicy 200 W, o pierwszych godzinach z systemem Zephyr RTOS oraz o kompensacji zasilacza według metodyki Christophe’a Basso. Repetytorium o oscyloskopach wkracza w drugą część, a osobny cykl przybliży wzmacniacze operacyjne i komparatory. Nie zabraknie też nowych podzespołów, modułów i krótkich prezentacji – bo nawet przy rosnącej liczbie materiałów tutorialowych chcemy, by EP pozostała oknem na to, co właśnie pojawia się na rynku. Życzę inspirującej lektury i – jak zawsze – czekam na Państwa listy oraz uwagi.

Wiesław Marciniak

Kilka refleksji starego prenumeratora

Jestem prenumeratorem jeszcze z czasów „Radioelektronika” z lat osiemdziesiątych. Piszę to trochę z uśmiechem, bo od tamtego czasu zmieniło się niemal wszystko: technika, rynek, sposób projektowania, dostęp do podzespołów, narzędzia, a nawet sam sposób czytania czasopism. A jednak potrzeba dobrej, rzetelnej, inspirującej prasy technicznej pozostała taka sama. Z dużym zainteresowaniem obserwuję ostatnie zmiany w „Elektronice Praktycznej”. Wcześniej podobną przemianę zauważyłem w „Elektronice dla Wszystkich”. Z pisma, które przez pewien czas kojarzyło mi się raczej z prostymi robótkami ręcznymi, ponownie stało się ono ciekawym wydawnictwem, po które chętnie sięgam, szukając projektów, inspiracji i pomysłów dydaktycznych. Mam też wgląd w wydawnictwa zagraniczne, które regularnie przeglądam, dlatego tym bardziej doceniam próbę odświeżenia polskich tytułów technicznych. Wracając jednak do „Elektroniki Praktycznej” – zmiana jest bardzo interesująca. Zastanawiam się jedynie, jak zostanie zagospodarowana przestrzeń pomiędzy „Elektroniką Praktyczną”, „Elektronikiem” a „Elektroniką dla Wszystkich”. „Elektronik” jest pismem wyraźnie profesjonalnym, branżowym, skierowanym do przemysłu i zawodowych elektroników. „EdW” ma naturalnie bardziej edukacyjny i hobbystyczny charakter. „EP” zawsze znajdowała się gdzieś pomiędzy: blisko praktyki, ale jednocześnie z ambicją pokazywania rozwiązań, które można wykorzystać także w poważniejszych zastosowaniach. I właśnie w tym miejscu widzę dla niej bardzo ciekawą przestrzeń.

Na obecnym etapie mogę podsumować kilka rzeczy:

Kursy uważam za bardzo potrzebne. Z kilku skorzystałem osobiście, z wielu korzystałem jako z materiałów pomocniczych w dydaktyce, ponieważ jestem wykładowcą. Dobre kursy techniczne mają ogromną wartość, szczególnie wtedy, gdy nie są tylko suchym przeglądem teorii, lecz prowadzą czytelnika od podstaw do praktycznego zastosowania. To jest coś, czego wciąż bardzo potrzeba – szczególnie dziś, kiedy młodzi ludzie często zaczynają od gotowych modułów, bibliotek i filmów w Internecie, ale nie zawsze rozumieją, co dzieje się „pod spodem”.

Projekty są drugim filarem pisma. Wielokrotnie korzystałem z Waszych opracowań – czasem bezpośrednio, budując układ zgodnie z opisem, częściej jednak wykorzystując pojedyncze rozwiązania, fragmenty koncepcji lub sposób podejścia do problemu we własnych aplikacjach. Mam świadomość, że projekty elektroniczne często krążą wokół kilku powtarzalnych obszarów: zasilaczy, sterowników, pomiarów, audio, komunikacji, automatyki, systemów mikroprocesorowych. Trudno za każdym razem wymyślić coś całkowicie nowego. Ale nawet inne spojrzenie na znany problem, zastosowanie niestandardowego algorytmu, ciekawej architektury układu albo nowego podzespołu może być bardzo wartościowe – zarówno edukacyjnie, jak i poznawczo.

Felietony to osobna sprawa. Dzisiejszy tekst – rewelacja. Wszyscy pewnie rzeczy wiemy, ale nie zawsze chcemy o nich głośno mówić. Dotyczy to szczególnie zależności od wielkich dostawców oprogramowania, usług chmurowych i systemów komunikacji. Nawet w mojej Alma Mater, po sugestiiach dotyczących większej niezależności cyfrowej, szybko okazało się, że „są umowy”, „musi być Microsyft” i nikt nie będzie eksperymentował z alternatywnymi rozwiązaniami. W praktyce oznacza to, że wyniki badań, dokumenty, dane dydaktyczne i wiele innych informacji trafiają do ekosystemów, nad którymi jako użytkownicy mamy bardzo ograniczoną kontrolę, większą ma usa wraz z kumplami. Można oczywiście powiedzieć: trudno, takie czasy. Ale może po ostatnich wybrykach wielkiego brata ktoś jednak zacznie się nad tym poważniej zastanawiać.

Ciekawy wydaje mi się także dział poświęcony modułom z Chin. Jeszcze nie mam o nim ostatecznego zdania, ale po dzisiejszym artykule zamówiłem już opisywaną kamerę do testów. Warto wiedzieć, co faktycznie dzieje się w tych rozwiązaniach, jak są zbudowane, jakie mają możliwości i ograniczenia. Pozostaje oczywiście problem czasu – bo kiedy moduł przyjdzie, zapewne kosztem innych aktywności będę go sprawdzał. Z drugiej strony, trudno nie zauważyć, że w „państwie środka” dzieje się dziś bardzo dużo. Czasem mam wrażenie, że więcej niż po naszej zachodniej stronie, zwłaszcza tej dalszej. Nowe moduły, układy, czujniki, kamery, radiomodemy, sterowniki, przetwornice i kompletne platformy pojawiają się w tempie, za którym trudno nadążyć.

Tematyka audio zawsze będzie na topie. I dobrze. Audio ma w sobie coś wyjątkowego: łączy elektronikę, akustykę, psychoakustykę, estetykę, pomiary i subiektywne wrażenia

użytkownika. Można pisać o lampach, układach dyskretnych, wzmacniaczach operacyjnych, DSP, aktywnych zwrotnicach, korekcji pomieszczeń, wzmacniaczach klasy D i profesjonalnych systemach nagłośnieniowych. Oczywiście tutaj pojawia się pewne pokrycie tematyczne z „EdW”, ale myślę, że da się to sensownie pogodzić. „EdW” może bardziej prowadzić początkujących za rękę, natomiast „EP” może pokazywać rozwiązania dojrzałe, bardziej pomiarowe, systemowe i aplikacyjne.

Bardzo ważnym działem jest również metrologia. W tej tematyce dzieje się dużo, a jednocześnie użytkownicy często są pozostawieni sami sobie. Recenzje sprzętu pomiarowego są potrzebne, ale dobrze byłoby, aby były pisane przez rzeczywistych użytkowników, a nie wyłącznie przez firmy sprzedające aparaturę. Różnica jest zasadnicza. Użytkownik napisze, co działa dobrze, co przeszkadza, gdzie producent przesadził z obietnicami, a które funkcje okazują się naprawdę przydatne w codziennej pracy. Firma handlowa, nawet jeśli działa uczciwie, zawsze będzie miała pewien konflikt interesów. Dlatego rzetelne, praktyczne recenzje oscyloskopów, analizatorów, mierników, kamer termowizyjnych, obciążen elektronicznych czy analizatorów mocy byłyby bardzo cennym elementem pisma.

Za bardzo dobry pomysł uważam także przegląd nowych podzespołów. Sam nie jestem w stanie regularnie śledzić stron wielu producentów, not aplikacyjnych, premier układów scalonych i nowych modułów. A przecież często jeden ciekawy element potrafi uruchomić cały projekt. W tym sensie taki dział może być dla czytelnika nie tylko informacją, ale także inspiracją.

Z mojego – podkreślam: tylko mojego – punktu widzenia warto byłoby rozszerzyć tematykę technologii radiowych, zwłaszcza w kontekście autonomicznych systemów sterowania. Nie myślę tutaj wyłącznie o prostym DIY, choć i tam jest miejsce na ciekawe konstrukcje. Bardziej chodzi mi o szersze spojrzenie: aparatury RC, systemy ISM, LoRa, telemetryczne łącza danych, FPV, radiomodemy, oprogramowanie, konfigurację, odporność transmisji, bezpieczeństwo, zasięgi, opóźnienia, a także rozwiązania profesjonalne, półprofesjonalne i MIL, w tym takie, w których konfiguracja systemu ma kluczowe znaczenie. Rynek z oczywistych powodów bardzo się poszerzył, a zainteresowanie bezzatogowymi platformami, robotyką mobilną, systemami autonomicznymi i zdalnym sterowaniem będzie raczej rość niż malać.

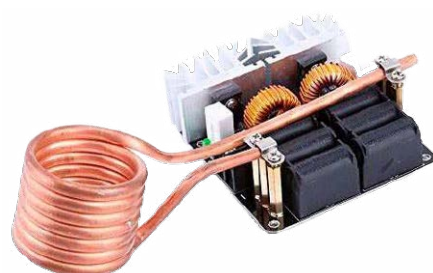
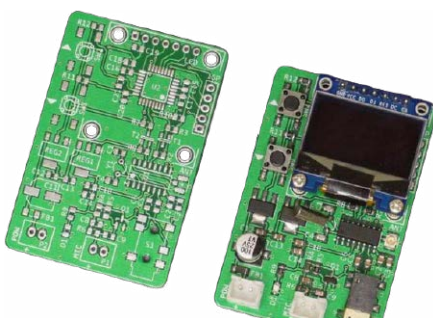
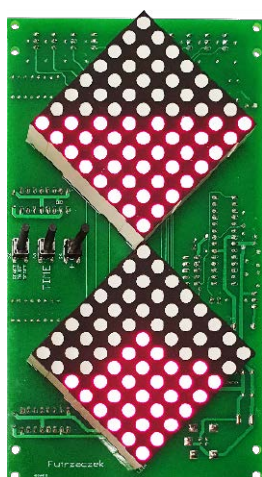
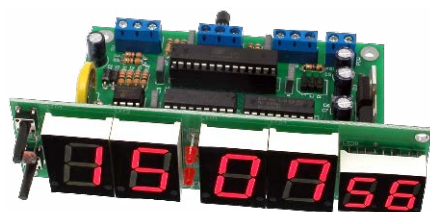
W dzisiejszym numerze pojawiły się również dwa wątki związane z energoelektroniką, którą sam zawodowo się zajmuję. Mam wrażenie, że w ostatnich latach tej tematyki było w „EP” stosunkowo mało. Tymczasem jest to obecnie jeden z najważniejszych obszarów elektroniki praktycznej. W połączeniu z OZE, szczególnie fotowoltaiką, magazynami energii, układami ładowania, przetwarzaniem mocy, mikrosieciami i systemami smart grid, otwiera się ogromna przestrzeń dla artykułów technicznych. Oczywiście pojawiają się tu niebezpieczne napięcia, znaczne prądy i zagrożenia bezpieczeństwa, ale właśnie dlatego warto o tym pisać dobrze, odpowiedzialnie i praktycznie. Dla mniej zaawansowanych może być „EdW”, natomiast „EP” mogłaby spokojnie wejść głębiej w przekształtniki, sterowanie silników, pomiary jakości energii, układy magazynowania, przetwornice DC/DC, inwertery, zabezpieczenia, algorytmy MPPT czy współpracę urządzeń z siecią.

Przypomina mi się przy tej okazji opublikowany niedługo na łamach „EP” jednofazowy falownik. Najpierw zbudowałem wersję zgodnie z projektem, później rozwijałem aplikację, a dziś podobne rozwiązania stosuję w wielu własnych układach sterowania silnikami indukcyjnymi. I to jest właśnie ogromna wartość dobrego projektu publikowanego w czasopiśmie technicznym. Nie musi on być od razu gotowym produktem przemysłowym. Czasem ważniejsze jest to, że pokazuje ideę, strukturę rozwiązania i sposób myślenia. Taki projekt może stać się punktem wyjścia do własnych eksperymentów, modyfikacji i zastosowań.

Dlatego mocno trzymam kciuki za pozytywny odbiór wprowadzonych zmian. Być może jestem już dinozaurem, ale nadal lubię po-trzymać w rękach prawdziwą gazetę, książkę czy inne wydawnictwo. Jednocześnie przyznaję, że pierwsze czytanie „EP” coraz częściej odbywa się u mnie na komputerze. Do e-wydań już się przyzwyczaiłem, choć papier wciąż ma dla mnie swój niepowtarzalny urok.

Życzę powodzenia, dobrych autorów, odważnych tematów i czytelników, którzy nie tylko czytają, ale także budują, mierzą, sprawdzają, psują i poprawiają. Bo chyba właśnie o to w elektronice praktycznej chodzi.

Do następnego numeru
Tomasz K.



NIE PRZEOCZ

Nowe podzespoły 6

Z PRAKTYKI EKSPERTÓW

AI w pracowni elektronika konstruktora..... 12
 Zephyr RTOS – pierwsze 48 godzin: od Blinky do BLE..... 24
 Półprzewodniki szerokoprzerwowe GaN vs. SiC
 – który wybrać do przetwornicy 200 W? 34
 Analiza, symulacja i eksperyment
 – klucz do udanego projektu kompensacji zasilacza 42

PODZESPOŁY

Przekaźniki półprzewodnikowe SSR
 – budowa, działanie i zastosowanie 56

PREZENTACJE

Joy-Car Calliope – robot, który uczy robotyki i programowania..... 59
 Jak rozwiązania półprzewodnikowe ewoluowały,
 by wspierać przemysł kosmiczny: od Apollo 11 do Starlink..... 92
 Shelly: ochrona przed zalaniem i zamek bez klucza 127

W PRACOWNI

Cyfrowy nadajnik FM 60
 Nagrzewnice indukcyjne ZVS 68
 Zegary hobbystyczne 78

KITY AVT

Elektroniczna klepsydra LED..... 86

MODUŁY Z CHIN

M5Stack Cardputer ADV z modułem LoRa-1262.
 Kieszonkowy komputer (nie tylko) edukacyjny na ESP32-S3..... 96

AI W ELEKTRONICE

AI wbudowane we współczesne urządzenia elektroniczne.
 Sprzęt, który musiał dorosnąć 100

AUDIO BEZ TAJEMNIC

Synteзаторы dźwięku, część 8. Filtry sterowane napięciem (2)..... 104

REPETYTORIUM

Oscyloskopy 2026 – od lampy Brauna
 do 12 bitów na biurku konstruktora, część 2 106
 Wzmocniacze operacyjne i komparatory 119

KURS

Kurs RFID z PN532, rozdział 1.
 RFID od zera: podłączamy PN532 i piszemy pierwszą klasę 128

FELIETON

„Jaki mikrokontroler wybrać do ...”, czyli przypowieść o Wojtku 138

Prenumerata 2
 Od wydawcy 3
 Z poczty 4

Nowe podzespoły

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



Moduł mikrowyświetlacza OP02220

OP02220 jest mikrowyświetlaczem typu LCOS (*Liquid Crystal on Silicon*). W strukturze modułu dostępny jest sterownik oraz bufor pamięci. Przekątna matrycy wynosi 0,39". Obsługiwane są następujące rozdzielczości obrazu: 1080p (przy 60 kl./s) oraz 720p (przy 120 kl./s). Transmisja danych realizowana jest czterema liniami interfejsu MIPI-DSI. Informacja o kolejnych pikselach przesyłana jest w formacie RGB888, w którym każda barwa składowa (czerwona, zielona i niebieska) kodowana jest na 8 bitach.

W podzespole dostępne są zarówno funkcje odwracania obrazu w pionie lub poziomie, jak i dedykowany blok adresacji danych do matrycy mikrowyświetlacza. Moduł OP02220 przeznaczony jest do zastosowań m.in. w aplikacjach rozszerzonej rzeczywistości (AR), projektorach PICO oraz wyświetlaczach przeziernych (*Head-up Display*). Wymiary modułu wynoszą: 25,7×12,6×3,33 mm. Przewidziana jest w nim możliwość przesunięcia obrazu o ±16 pikseli celem przeprowadzenia procesu kalibracji. W trakcie pracy pobór mocy równa się 330 mW. Natomiast w trybie czuwania wartość poboru nie przekracza 10 mW, a zakres temperatury pracy mieści się w przedziale: od -10 do 70°C.

www.ovt.com

Płaskie superkondensatory z serii SCCA

Seria superkondensatorów SCCA obejmuje podzespoły bierne w obudowie przypominającej monetę (typu *coin-cell*), które charakteryzuje napięcie znamionowe 5,5 V. Zakres pojemności podzespołów wynosi: od 100 mF do 1500 mF. Elementy przystosowane są do montażu przewlekane (THT) na płytkach drukowanych PCB. Niewielkie wymiary ułatwiają stosowanie



superkondensatorów SCCA w rozwiązaniach o ograniczonej przestrzeni montażowej.

Najważniejsze cechy serii SCCA: to niski prąd upływu, duża gęstość mocy oraz wysoka liczba cykli ładowania i rozładowania. Zakres temperatury pracy superkondensatorów mieści się w przedziale: od -25 do 70°C. Z kolei temperatura przechowywania wynosi: 0...60°C, przy wilgotności względnej do 70% RH.

We wszystkich wariantach wyprowadzenia są cynowane i przystosowane do lutowania falowego. Dobór wariantu powinien uwzględniać m.in. rodzaj obciążenia, maksymalne napięcie pracy i wartość szeregowej rezystancji zastępczej (ESR). Równocześnie należy zapewnić spełnienie wymagań czasowych i prądowych aplikacji w granicach specyfikacji projektowych.

Superkondensatory z serii SCCA są zgodne z dyrektywami unijnymi RoHS i REACH. Zastosowania podzespołów obejmują w szczególności układy wymagające krótkotrwałego podtrzymania napięcia. Typowe aplikacje to dla przykładu: podtrzymanie pracy zegara czasu rzeczywistego (RTC), zapobieganie utracie danych w pamięciach podczas zaniku pamięci oraz wsparcie dla obciążeń w układach sterowania oraz poszczególnych obwodach elektronicznych.

www.schurter.com

IQXC-26 Miniature Quartz Crystal



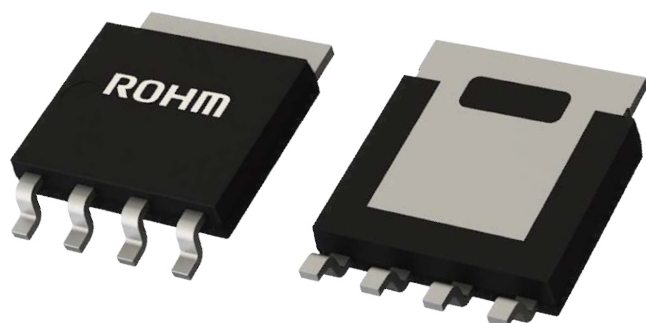
Rezonator kwarcowy IQXC-26

Rezonator IQXC-26 dostępny jest w szczelnej obudowie ceramicznej o wymiarach: 1,6×1,2×0,4 mm. Element chroniony jest przed wpływem zakłóceń elektromagnetycznych (EMI) i czynników środowiskowych. Charakteryzuje go również częstotliwość pracy: 24...60 MHz. W zależności od wariantu tolerancja częstotliwości wynosi: od ±10 ppm do ±50 ppm. Stabilność w funkcji temperatury mieści się

w przedziale: od ± 15 ppm do ± 50 ppm. Natomiast dryft starzeniowy nie przekracza ± 3 ppm rocznie, a temperatura pracy jest równa: od -20 do 70°C lub od -40 do 85°C .

Pojemność obciążenia CL w rezonatorze IQXC-26 zawiera się w zakresie: 8...30 pF. Pojemność równoległa Co to z kolei wartość nie wyższa niż 3 pF. Dodatkowo poziom wzbudzenia (drive level) osiąga 100 μW . Spełniane są przy tym wymagania normy MIL-STD-202. Zakres temperatury przechowywania podzespołu wynosi: od -55 do $+125^{\circ}\text{C}$. Rezonator jest zgodny z dyrektywami unijnymi RoHS i REACH. Stosowany jest on jest m.in. w torach taktowania mikroprocesorów. Spotyka się go także w mikrokontrolerach oraz systemach komunikacji radiowej, szczególnie tych w których przeprowadzana jest zaawansowane przetwarzanie sygnałów (DSP).

www.iqdfrequencyproducts.com



Tranzystor mocy AGO40FGS4FRA

AGO40FGS4FRA to tranzystor krzemowy MOSFET z kanałem wzbogacanym typu N. Element dyskretny, który charakteryzuje się: napięciem maksymalnym „dren-źródło” VDSS równym 40 V oraz prądem drenu ID o wartości 120 A. Jest to wartość obowiązująca przy temperaturze obudowy TC wynoszącej 25°C . Podzespół przeznaczony jest do zastosowań przemysłowych lub motoryzacyjnych. Względem przyrządu półprzewodnikowego zachowuje się zgodnie ze standardem AEC-Q101. Tranzystor AGO40FGS4FRA oferowany jest w obudowie o wymiarach: $6 \times 4,9 \times 1,2$ mm. Obudowa przystosowana jest do montażu powierzchniowego (SMT) a sam element może zostać użyty na wszystkich płytkach PCB.

W stanie włączenia wartości rezystancji RDS(on) wynoszą: 1,8 m Ω – przy napięciu „bramka-źródło” VGS = 4,5 V oraz 1,1 m Ω – przy VGS = 10 V. Dzięki redukcji całkowitego ładunku bramki Qg, istnieje możliwość stosowania tranzystora AGO40FGS4FRA w aplikacjach o wysokiej częstotliwości przełączania. Wśród aplikacji wymienić można m.in.: przetwornice DC/DC, falowniki i układy napędowe w pojazdach elektrycznych (EV). Dodatkowo moc rozpraszana PD nie przekracza 150 W pod warunkiem zagwarantowania właściwych warunków chłodzenia.

Zakres temperatury pracy AGO40FGS4FRA mieści się w przedziale: od -55 do 175°C . Wszystkie przytoczone

parametry predestynują tranzystor do zastosowań przede wszystkim w układach zasilania, sterownikach silników oraz systemach energoelektronicznych o obciążeniach często zmieniających się w czasie.

www.rohm.com

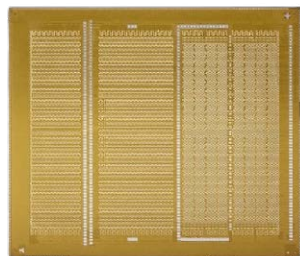


Przełącznik mikrofalowy PE42448

Firma pSemi wprowadza na rynek przełącznik mikrofalowy typu SPDT (Single Pole Double Throw) o symbolu PE42448 wykonany w technologii UltraCMOS, przeznaczony do pracy w paśmie częstotliwości: od 10 MHz do 6 GHz. Podzespół zaprojektowano z myślą o aplikacjach, w których znaczenie ma ograniczenie strat mocy wraz z eliminacją przesłuchów międzykanałowych. Parametry pasożytnicze w strukturze PE42448, w tym pojemności i rezystancje, zostały całkowicie zredukowane. Jest to rozwiązanie przeznaczone do użycia zwłaszcza w szynach fazowanych i stacjach bazowych 5G oraz wszędzie tam, gdzie wymagane jest szybkie i niezawodne przełączanie sygnałów mikrofalowych. Dzieje się to bez pogarszania parametrów z nimi związanych oraz bez wprowadzania zniekształceń na wysokim poziomie.

W systemach wielokanałowych przełącznik mikrofalowy PE42448 umożliwia praktyczną implementację torów odbiorczych i nadawczych. W ich kontekście należy pamiętać o zapewnieniu dopasowania impedancyjnego 50 Ω , dzięki czemu nie następuje przyrost wartości współczynnika fali stojącej (WFS), a przez to nie dochodzi do uszkodzenia podłączonych urządzeń. Prócz tego przewidziana jest możliwość użycia przełącznika w systemach o wysokiej skali integracji podzespołów. Przełącznik PE42448 rozszerza ofertę aplikacji radiowych firmy pSemi w zakresie techniki mikrofalowej. Jest on polecany ze względu m.in. na wysoką liniowość charakterystyk, niskie straty wtrąceniowe oraz niewygórowany zakres częstotliwości odpowiedni w szczególności dla mniej złożonych systemów telekomunikacji, a także radarów zarówno cywilnych, jak i wojskowych.

psemi.com



Układy pamięci HBM4

Układy HBM4 oparto na dwunastowarstwowym stosie matryc DRAM umożliwiającym uzyskanie pojemności: od 24 do 36 GB. Co więcej, dostępny jest również wariant o stosie szesnastowarstwowym, którego pojemność nie przekracza 48 GB. Podzespół wytwarzany jest w 2 procesach technologicznych: 10 nm – dla matryc DRAM oraz 4 nm – dla układów sterowania. Charakteryzuje go przepustowość maksymalna 3,3 TB/s, co w porównaniu do rozwiązań starszej generacji warunkuje przyrost wydajności około 2,7 razy.

W celu obniżenia strat mocy zastosowane zostały rozwiązania redukujące pobór mocy. Równocześnie dokonano optymalizacji ścieżek, które doprowadzają napięcie zasilania do podzespołów w układach HBM4. W ten sposób podwyższono sprawność energetyczną układów, zwiększono ich zdolność do odprowadzania ciepła i sprawiono, że podzespoły mogą pracować w rozszerzonym zakresie temperatury otoczenia.

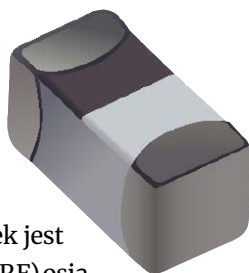
Rozwój rodziny układów obejmuje kolejne warianty, w tym HBM4E, dostosowane do specyficznych wymagań systemów obliczeniowych. Zwiększona przepustowość danych, przy wzroście sprawności, sprzyja ograniczeniu kosztów eksploatacyjnych, szczególnie wtedy, gdy układy HBM4 znajdują zastosowania w kartach graficznych lub w modułach, których najważniejszym przeznaczeniem są centra przetwarzania danych.

news.samsung.com

Wielowarstwowe cewki z serii CE0603G

Seria CE0603G obejmuje wielowarstwowe cewki przeznaczone do montażu powierzchniowego (SMT) dostępne w obudowie o wymiarach: 0,6×0,3×0,3 mm. Najważniejszym z parametrów cewek jest częstotliwość rezonansu własnego (SRF) osiągająca maksymalnie 18 GHz, która maleje wraz ze wzrostem indukcyjności w zakresie: 0,3...39 nH. Przytoczony zakres jest zbiorczy, a każdą cewkę z serii CE0603G charakteryzuje jedna z jego wartości.

Najwyższy prąd znamionowy mieści się w przedziale: 120...850 mA, przy rezystancji: 0,07...2,4 Ω i tempera-



turze pracy: od -55 do 125°C. Klasa wrażliwości na wilgoć odpowiada poziomowi MSL 1, co eliminuje konieczność zapewniania dedykowanych warunków przechowywania podzespołów przed ich wykorzystaniem. Oprócz tego obudowa cewek wykonana jest ze specjalnego materiału ceramicznego, pozwalającego na użycie elementów w standardowych warunkach przemysłowych.

Końcówki cewek pokryto warstwami Ag/Ni/Sn przystosowanymi do lutowania rozpliwowego. Zalecany profil lutowania dopuszcza temperaturę szczytową do 260°C. Elementy serii CE0603G są zgodne z dyrektywą unijną RoHS. Zastosowania cewek uwzględniają m.in. filtry przeciwzakłóceniu, niskonapięciowe układy zasilania oraz wzmacniacze i nadajniki radiowe, w przypadku których wymagana jest stabilna wartość parametrów pracy, przy wysokiej częstotliwości sygnałów emitowanych oraz ograniczonych stratach mocy.

bourns.com

Lokalizator AirTag najnowszej generacji

Najnowsza wersja lokalizatora AirTag stanowi rozwinięcie poprzedniej wersji urządzenia przeznaczonego do identyfikacji oraz określania położenia przedmiotów w ramach systemu Find My. W konstrukcji lokalizatora zastosowano drugą generację układu Ultra Wideband (UWB), która zwiększa zasięg precyzyjnego określania pozycji względem wersji wcześniejszej. Prócz tego udoskonalono interfejs Bluetooth oraz zwiększono poziom ciśnienia akustycznego, które generowane jest przez wbudowany przetwornik elektroakustyczny. Emitowany przez przetwornik sygnał jest teraz głośniejszy o 50%, dzięki czemu odkrycie AirTag stało się dużo łatwiejsze zwłaszcza w miejscach o podwyższonym poziomie hałasu.

Zintegrowana funkcja Precision Finding wykorzystuje dane z interfejsu UWB i następujące sprzężenia zwrotne: wizualne, dźwiękowe oraz haptyczne, w celu wskazania kierunku i odległości do lokalizatora. Po aktualizacji systemu watchOS do wersji 26.2.1 możliwe jest inicjowanie procedury precyzyjnego lokalizowania AirTag z poziomu zegarków Apple Watch serii 9 lub nowszej oraz modeli Ultra 2, bez konieczności korzystania do tego celu ze smartfonów iPhone.

Emitowane przez AirTag sygnały Bluetooth są anonimowo odbierane przez kompatybilne urządzenia znajdujące się w pobliżu, a następnie przekazywane do systemu



lokalizacyjnego. Pozwala to określić przybliżoną pozycję urządzenia w sytuacji, kiedy lokalizator znajduje się poza bezpośrednim zasięgiem właściciela.

Wymiary AirTag nie uległy zmianie. Zachowano kompatybilność urządzenia z dotychczasowymi akcesoriami. Zasilanie podzespoły realizowane jest z wymiennej baterii litowej typu CR2032. Parametry zasilania i integracja z ekosystemem urządzeń pozostają zgodne, a nowa wersja lokalizatora została wprowadzona do sprzedaży bez zmiany sugerowanej ceny detalicznej.

www.apple.com



Kondensatory elektrolityczne aluminiowe z serii AEW

Kondensatory serii AEW przeznaczone są do pracy w szerokiej gamie obwodów zasilania, wszędzie tam, gdzie przestrzeń montażowa jest ograniczona. W konstrukcji podzespołów zastosowano anodowaną folię aluminiową o zwiększonej pojemności wraz z okładkami, które cechuje podwyższona zdolność do zgromadzania ładunków.

W połączeniu z dedykowanym elektrolitem oznacza to możliwość pracy w wysokiej temperaturze otoczenia (w zakresie: od -40 do 105°C) oraz obniżoną wartość szeregową rezystancji zastępczej (ESR).

Oferowany zakres napięć roboczych obejmuje wartości: $400\ldots 450\text{ V DC}$. Natomiast pojemności znamionowe zawierają się w przedziale: $82\ldots 300\text{ }\mu\text{F}$. Dodatkowo trwałość eksploatacyjna wynosi co najmniej 3000 godzin.

W porównaniu z serią HXW odnotowuje się wzrost pojemności do prawie 20%, przy zmniejszeniu objętości obudowy do 17%. Umożliwia to realizację bardziej zwartej topologii filtrów wejściowych i innych obwodów w zasilaczach impulsowych.

Kondensatory dostępne są w obudowach o wymiarach: od $\varnothing 1,6 \times 2,5\text{ cm}$ do $\varnothing 1,8 \times 5\text{ cm}$. Wymiary odpowiadają typowym rozstawom otworów montażowych stosowanych w płytkach PCB dla elementów osiowych. Oprócz tego parametry elektryczne predysponują serię AEW do zastosowań w przemysłowych układach zasilania oraz zasilaczach serwerowych i adapterach sieciowych wymagających skutecznego filtrowania przy podwyższonym obciążeniu na ich wejściach.

www.rubycon.co.jp



Dioda laserowa UV o długości fali 379 nm i mocy optycznej 1 W

Oferta Nuvoton Technology została rozszerzona o półprzewodnikową diodę laserową emitującą promieniowanie ultrafioletowe (UV) o długości fali 379 nm. W trybie emisji ciągłej (continuous wave) podzespół charakteryzuje moc optyczna 1 W, która obowiązuje przy temperaturze obudowy 25°C . Element umieszczony jest w metalowej obudowie TO-9 wykonanej z materiałów o wysokiej przewodności cieplnej. Umożliwia ona efektywne usuwanie

ciepła, co stabilizuje parametry pracy i w rezultacie przyczynia się do wydłużenia czasu eksploatacji podzespołu.

Struktura diody złożona jest ze specjalnie dobranych warstw półprzewodnikowych. Równocześnie ograniczenie strat absorpcji i napięcia pracy skutkuje wzrostem sprawności konwersji energii elektrycznej na promieniowanie UV. Parametry emisji wiązki predysponują diodę do zastosowań szczególnie w litografii bezmaskowej, w której laser prowadzony jest bezpośrednio po materiale światłoczułym.

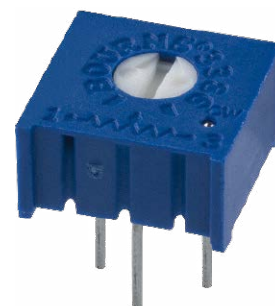
Dioda stanowi element serii źródeł UV, które stanowią zamienniki dla lamp rtęciowych. Jej aplikacje obejmują również: utwardzanie materiałów światłoutwardzalnych, znakowanie, druk 3D oraz aplikacje biomedyczne.

www.nuvoton.com

Potencjometr cermetowy 3386P-1-504LF

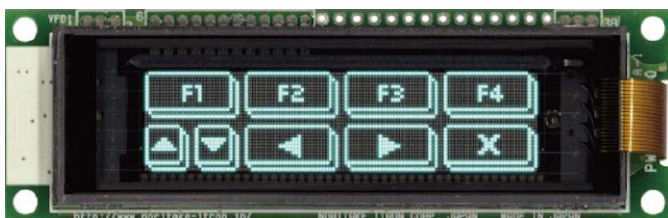
3386P-1-504LF jest potencjometrem dostrojczym (trymerem) o elemencie rezystancyjnym, który wykonany jest z cermetu (inaczej: cermetalu, spieku ceramiczno-metalowego). Rezystancja znamionowa potencjometru wynosi 500 k Ω ,

przy wartości tolerancji $\pm 10\%$. Temperaturowy współczynnik rezystancji nie przekracza $\pm 100\text{ ppm}/^{\circ}\text{C}$ i został określony dla przedziału temperatury: od -55 do 125°C . Przedział dotyczy temperatury roboczej, w której potencjometr pracuje przy zachowaniu deklarowanych względem niego parametrów elektrycznych.



Potencjometr 3386P-1-504LF ma uwzględnioną regulację na górze, która pozwala na dostosowanie elementu do obwodów płytek PCB. Zastosowana w podzespolu obudowa ma wymiary: 9,53×9,53 mm, zatem jej podstawa to kwadrat. Obudowa wykonana jest z tworzywa sztucznego. Jest ona dostępna w kolorze niebieskim. Oprócz tego potencjometr przeznaczony jest do montażu przewlekane (THT). Element wyposażony jest w wyprowadzenia, które można umieszczać zwłaszcza w otworach metalizowanych płytek, gdzie są dołączane w celu zapewnienia trwałych połączeń z istniejącymi obwodami elektrycznymi.

bourns.com



Wyświetlacze fluorescencyjne z serii GU-D

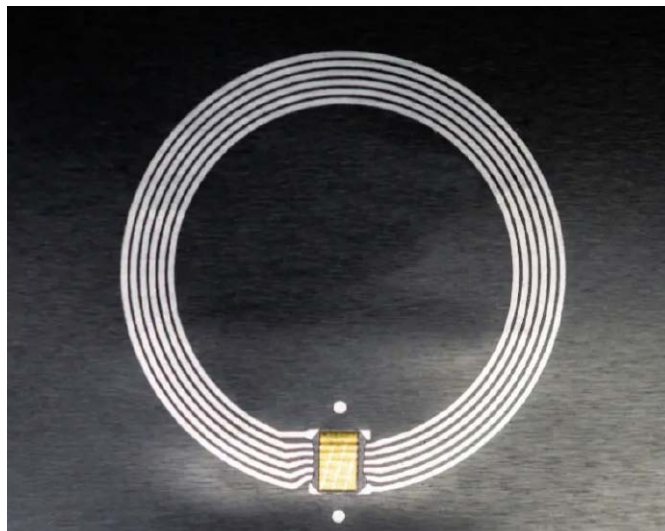
Seria GU-D obejmuje monochromatyczne wyświetlacze VFD (Vacuum Fluorescent Display) charakteryzowane przez wysokie kąty widzenia, emitujące charakterystyczne światło w kolorze niebiesko-zielonym. Wyświetlacze serii pracują w temperaturze otoczenia: od -40 do 85°C i mogą być używane w szerokiej gamie aplikacji, w tym w rozwiązaniach działających w wymagających warunkach środowiskowych, takich jak: systemy przemysłowe, układy automatyki lub urządzenia pomiarowe.

Komunikacja z wyświetlaczami realizowana jest za pośrednictwem interfejsów: I²C, SPI lub UART. Interfejsy umożliwiają sterowanie wyświetlaczem przy użyciu dedykowanego zestawu poleceń, pozwalających m.in. na wyświetlanie tekstu i grafik. Przewidziane są również cztery linie ogólnego przeznaczenia (GPIO), które można podłączyć do urządzeń zewnętrznych. Zintegrowany mikrokontroler umożliwia obsługę danych z interfejsów i linii, co pozwala na dynamiczne sterowanie informacjami prezentowanymi na wyświetlaczach, zgodnie z tym, co zakładają użytkownicy.

Dodatkowo wyświetlacze z serii GU-D mogą zawierać pojemnościowe panele dotykowe. Panele wykonane są w technologii cienkowarstwowego przewodnika o niskiej rezystancji, co zwiększa stosunek sygnału do szumu oraz rozszerza margines detekcji dotyku. Niektóre z paneli wykorzystują metodę pomiaru pojemności wzajemnej dla sterowania wyświetlaczami. Zastosowanie metody umożliwia poprawną detekcję dotyku zwłaszcza w obecności wilgoci lub gdy ręce są w rękawicach ochronnych. Dzięki temu, interfejsu użytkowników pozostają niezawodne

w wymagających środowiskach pracy, czyli wyświetlacze GU-D warunkują tworzenie nowoczesnych i trwałych rozwiązań łączących zalety technologii VFD z bezawaryjną obsługą dotykową – przy zachowaniu odporności na czynniki środowiskowe.

www.noritake-elec.com



Znacznik identyfikacyjny NFC PR1301

PR1301 jest pasywnym znacznikiem identyfikacji radiowej, przeznaczonym do komunikacji bliskiego zasięgu. Transmisja danych odbywa się na częstotliwości 13,56 MHz, w trybie półdupleksowym, zgodnie ze standardem NFC Forum Type 5. Parametry komunikacyjne znacznika spełniają wymagania normy ISO 15693. W strukturze elementu „zaszyty” jest unikatowy identyfikator UID oraz uwzględniona jest pamięć użytkownika o pojemności 96 B, która dostępna jest z poziomu interfejsu NFC.

Przepustowość transmisji danych wspierana przez znacznik PR1301 nie przekracza 26,48 kb/s. Rozwiązanie nie wymaga zasilania zewnętrznego, ponieważ energia pobierana jest z pola elektromagnetycznego, generowanego przez czytnik NFC.

Element współpracuje z czytnikami przemysłowymi oraz urządzeniami mobilnymi wyposażonymi w interfejs NFC, w tym z urządzeniami pracującymi pod kontrolą systemów operacyjnych: Android lub iOS. Jest to lekki i równocześnie giętki komponent przeznaczony do użycia wraz z wkładkami RFID wykonanymi z papieru lub tworzyw sztucznych. Dzięki temu, znacznik PR1301 może być stosowany m.in. w etykietach produktów oraz na ich opakowaniach, nie powodując istotnego zwiększenia ich grubości ani sztywności. Zatem to rozwiązanie dla nowoczesnych systemów znakowania, które może być integrowane z różnymi aplikacjami i wykorzystywane również do identyfikacji produktów, pod warunkiem zastosowania technologii NFC zgodnej ze specyfikacją NFC Forum Type 5.

www.pragmaticsemi.com



Detektor podczerwieni PCI-3TE-12-1x1-TO8-wZnSeAR-36 oparty na heterostrukturach HgCdTe

Detektor PCI-3TE-12-1x1-TO8-wZnSeAR-36 chłodzony jest trójstopniowym modułem termoelektrycznym, co zapewnia stabilne warunki pracy oraz powtarzalność parametrów metrologicznych. Rozwiązanie przeznaczone jest do detekcji promieniowania w zakresie podczerwieni (IR). Charakterystyczna długość fali detektora (λ_{spec}) wynosi 12 μm . Odpowiedź widmowa rozciąga się do długości $\lambda_{\text{cut-off}} = 14 \mu\text{m}$, a maksimum czułości przypada w pobliżu długości fali równej 10,6 μm .

Powierzchnia czynna elementu światłoczułego w detektorze (Ao) wynosi: 1x1 mm. Detektor cechuje się kątem akceptacji promieniowania równym 36°. Zintegrowane okno optyczne wykonane jest z selenku cynku (ZnSe) i pokryte je powłoką przeciwoodbiciową. Jego konstrukcja okna obejmuje klin, który przeciwdziała powstawaniu interferencji optycznych. Typowa stała czasowa detektora nie przekracza 5 ns. Dla długości charakterystycznej fali najmniejsza czułość prądowa wynosi 0,07 A/W. Natomiast nominalna rezystancja elementu światłoczułego sięga około 200 Ω .

Napięcie polaryzacji detektora (V_b) nie może przekraczać 0,9 V. Dodatkowo w zakresie niskich częstotliwości obserwuje się pogorszenie stosunku sygnału do szumu wynikające z obecności szumu typu 1/f. Detektor

PCI-3TE-12-1x1-TO8-wZnSeAR-36 zamknięto w obudowie TO-8 przystosowanej do montażu w szerokiej gamie układów pomiarowych. Konstrukcja obudowy zapewnia stabilne warunki pracy elementu oraz umożliwia integrację detektora z układami optycznymi szeroko stosowanymi w precyzyjnych systemach pomiarowych i analitycznych. vigophotonics.com

Czujnik ciśnienia MLX90809

Czujnik MLX90809 umożliwia precyzyjny pomiar ciśnienia względnego w różniowanych warunkach pracy. Konstrukcja elementu zapewnia podwyższoną odporność na zakłócenia elektromagnetyczne, a sam podzespół jest fabrycznie skalibrowany. Jest w nim stosowana kompensacja temperaturowa wraz z korekcją nieliniowości, co pozwala na utrzymanie stabilnych parametrów pomiarowych w całym zakresie temperatury pracy.

Niepewność pomiarowa czujnika wynosi $\pm 1,5\%$ zakresu pomiarowego. Parametry metrologiczne MLX90809 są utrzymywane w zakresie temperatury: od -40 do 150°C . Spełniane są również wymagania stawiane elementom stosowanym w aplikacjach motoryzacyjnych. Czujnik oferuje funkcje diagnostyczne, przeznaczone do detekcji uszkodzeń lub nieprawidłowości w działaniu układu.

W czujniku uwzględniona jest obudowa o podwyższonej wytrzymałości mechanicznej. Wraz z dedykowanym sposobem montażu ogranicza to dryft parametrów czujnika w trakcie eksploatacji. Dzięki temu, podzespół jest przystosowany do pracy w wymagających warunkach, typowych w szczególności dla szerokiej klasy rozwiązań przemysłowych, takich jak: układy sterowania silnikami, systemy pneumatyczne, czy instalacje hydrauliczne.

www.melexis.com



REKLAMA

Mnóstwo doskonałych projektów, tylko na:

EP.com.pl



AI w pracowni elektronika konstruktora

Stan sporu na połowę 2026 roku – co naprawdę działa, a co tylko obiecują filmiki.

Pół roku w tej dziedzinie to epoka. Kiedy zaczynał się 2025 rok, pod reklamami kursów programowania mikrokontrolerów trwał spór: „po co kursy, skoro AI wszystko pisze” kontra „AI konfabuluje, to zabawka”. Dziś, w połowie 2026, brzmi to jak spór o coś, co już się rozstrzygnęło – tyle że nie po żadnej ze stron.

Część dawnych zarzutów sceptyków po prostu się zdezaktualizowała. Narzędzia, które jeszcze w 2025 roku „raz działały, raz nie”, dziś w większości działają. Ale w tym samym czasie pojawił się nowy, twardo udokumentowany powód do ostrożności – i nie dotyczy on już tego, czy AI potrafi napisać działający kod, lecz tego, co ten kod jest wart, gdy spojrzeć na jego bezpieczeństwo, jakość w dłuższym horyzoncie i realny wpływ na pracę całego zespołu. Ten artykuł pokazuje, jak przez ostatni rok przesunęło się ognisko sporu i co z tego wynika konkretnie dla pracowni elektronika konstruktora – bo embedded, jak zobaczymy, gra według nieco innych reguł niż reszta świata oprogramowania.

1. Spór, który zdążył się zestarzeć

Żeby zrozumieć, gdzie jesteśmy w połowie 2026 roku, trzeba zobaczyć, jak szybko grunt usunął się spod nóg obu obozom.

Najpierw obozowi sceptyków, bo jego sztandary dowód zestarzał się najbardziej. W lipcu 2025 głośnym echem odbiło się randomizowane badanie organizacji METR: doświadczeni programiści, pracujący nad własnymi, dobrze znanymi projektami, byli z narzędziami AI powolniejsi – zadania zajmowały im o 19% więcej czasu, choć po badaniu szacowali, że AI przyspieszyło ich o 20% [1]. Ten rozjazd między odczuciem a pomiarem stał się ulubionym argumentem sceptyków. Problem w tym, że gdy METR spróbował powtórzyć badanie, ono się rozpadło – i to w sposób, który sam jest dowodem zmiany. Deweloperzy byli już tak przyzwyczajeni do AI, że odmawiali pracy bez niego, co uniemożliwiło utrzymanie grupy kontrolnej [2]. W aktualizacji z lutego 2026 METR napisał

wprost, że deweloperzy są na początku 2026 roku prawdopodobnie bardziej przyspieszani przez AI niż w szacunkach sprzed roku [2]; surowe dane pokazują już sygnał przyspieszenia. Zdanie „AI spowalnia pracę programistów” było więc prawdziwe dla narzędzi z początku 2025 – i nie należy go dziś cytować jako stanu obecnego.

Co napędza tę zmianę, widać nie w hasłach marketingu, lecz w tym, jak długie zadania modele potrafią dziś doprowadzić do końca. METR mierzy to wprost: sprawdza, jak rozbudowane zadanie programistyczne model wykonuje skutecznie w co najmniej połowie prób – bo im dłuższą i bardziej złożoną pracę ogarnia samodzielnie, tym jest zdolniejszy. Postęp w ciągu roku był uderzający. Jeszcze niedawno modele radziły sobie z zadaniami, które doświadczonemu programiście zajmują około godziny; na początku 2026 roku – z takimi, które zajęłyby mu kilkanaście godzin, a ta zdolność podwaja się mniej więcej co cztery miesiące [3]. To realna miara, dlatego praktycy zmienili zdanie. Simon Willison (współtwórca frameworka Django), wskazuje konkretny przełom: listopad 2025, gdy agenci kodowania przeszli ze stanu „w większości działa” do „faktycznie działa” [4]. Tu obóz entuzjastów ma więc rację, której nie miał jeszcze rok wcześniej.

Ale dokładnie wtedy, gdy ucichł stary spór o zdolność, narodził się nowy – o jakość i bezpieczeństwo – i ten ma już za sobą dane, których w 2025 roku brakowało. „Szybciej i sprawniej” nie znaczy „bezpieczniej”. Pokazuje to badanie firmy Veracode, która przetestowała ponad sto modeli,

każąc im rozwiązać zestaw typowych zadań i sprawdzając, czy w kodzie pojawiają się klasyczne, powszechnie znane luki bezpieczeństwa – te, które od lat figuruje na branżowych listach najczęstszych błędów (jak dziury pozwalające wstrzyknąć do programu obcy kod czy wykraść dane). Wynik: 45% kodu generowanego przez AI zawierało taką lukę [5]. Co istotniejsze – odsetek rozwiązań wolnych od tych błędów utrzymywał się na poziomie około 55% przez cały okres 2025–2026, mimo że pod względem samego działania kod stał się wyraźnie lepszy [6]. Mówiąc prościej: nowsze modele coraz częściej piszą kod, który robi to, co trzeba, ale nie piszą go bezpieczniej – mniej więcej co drugi fragment wciąż zawiera znaną podatność.

A na poziomie całych zespołów pojawił się paradoks, który powinien ostudzić najgorętsze głowy: mimo że z AI korzysta już ponad dziesięciu na dziesięciu programistów, a około jedna czwarta kodu trafiającego do produktów jest dziś generowana przez AI, łączna wydajność organizacji wzrosła o nie więcej niż 10% [3]. Indywidualny zysk czasu pojedynczego programisty nie przekłada się więc wprost na szybszą pracę firmy – bo czas zaoszczędzony na pisaniu wraca w postaci przeglądów, poprawek i łatania luk.

Z tego wyłania się teza całego artykułu, inna niż wygodny „złoty środek”. W połowie 2026 roku pytanie „czy AI potrafi” jest w dużej mierze rozstrzygnięte – potrafi, i z miesiąca na miesiąc lepiej. Ognisko sporu przesunęło się gdzie indziej: ku pytaniom „jakim kosztem dla bezpieczeństwa i jakości”, „czy rozumiem to, co podpisuję” oraz „gdzie

ZDOLNOŚĆ KONTRA ODCZUCIE

Co naprawdę zmierzono w latach 2025-2026



Wniosek

„AI spowalnia ekspertów” było prawdą dla narzędzi z początku 2025 r. — nie dla dzisiejszych. Zdolność realnie rośnie, ale poczucie „idzie mi szybciej” bywa zawodne i wymaga pomiaru.

Źródło: METR — RCT 2025; aktualizacja II 2026; ankieta V 2026.

Elektronika Praktyczna

Rysunek 1. Zdolność kontra odczucie. W badaniu z 2025 r. doświadczeni programiści byli z AI realnie wolniejsi, choć czuli przyspieszenie – wynik dotyczył jednak narzędzi sprzed roku. W 2026 r. zdolność modeli wyraźnie wzrosła. Źródło: METR

konkretnie, na jakim etapie projektu i pod jakim nadzorem AI realnie pomaga”. I właśnie tu zaczyna się historia specyficzna dla konstruktora – bo świat sprzętu i embedded odpowiada na te pytania inaczej niż świat aplikacji webowych, którego dotyczy większość przytoczonych liczb.

2. Przykładowy eksperyment

Przytoczone wcześniej liczby opisują świat oprogramowania w ogóle. Żeby zobaczyć, co naprawdę znaczą w pracowni elektronika, najlepiej prześledzić jedno konkretne starcie człowieka z narzędziem. Jeden z polskich konstruktorów embedded – czyli specjalista od układów sterowanych mikrokontrolerem, „wmontowanym” w urządzenie – opisał eksperyment przeprowadzony na żywo (otrzymaliśmy w poczcie redakcyjnej). Posadził dwa wiodące narzędzia AI do programowania, tak zwani agenci kodowania (programy, które nie tylko podpowiadają fragmenty kodu, ale samodzielnie planują pracę, piszą pliki i uruchamiają polecenia), przed realistycznym, choć celowo niełatwym zadaniem na popularnym mikrokontrolerze STM32.

Zadanie brzmiało zwyczajnie: odczytać temperaturę i ciśnienie z czujnika BMP280, połączanego z mikrokontrolerem magistralą I²C i wysłać wynik do komputera przez UART. To robota, którą doświadczony konstruktor wykonuje niemal odruchowo – a która dla modelu językowego jest trudniejsza, niż się wydaje.

Kluczowy był jeden zabieg. Autor świadomie wybrał na tyle świeżą wersję bibliotek i narzędzi producenta mikrokontrolera, że ich kompletnego kodu nie ma jeszcze w publicznym obiegu w internecie. To istotne, bo model językowy uczy się na tym, co już gdzieś napisano; gdy materiału brakuje, musi zgadywać przez analogię do wersji starszych. Eksperyment był więc celowo ustawiony tak, by sprawdzić narzędzie nie na „przeciwczonym” przykładzie, lecz na czymś, czego nie mogło wcześniej „widzieć” – a to, jak się okaże, jest w embedded sytuacją typową, nie wyjątkową.

Dwa style, dwie szkoły myślenia. Najciekawsze było to, jak różnie zachowały się oba narzędzia. Pierwszy agent wystartował z rozmachem: uruchomił tryb planowania, sam zadał kilka trafnych pytań (jak połączyć czujnik, czy konfigurować układ w graficznym generatorze, czy ręcznie w kodzie, jak formatować dane), rozpoznał strukturę projektu, wychwycił nowe nazewnictwo w bibliotekach producenta, założył nawet plik z konwencjami projektu – i od razu chciał wygenerować pełną, gotową implementację. Drugi agent zachował się przeciwnie i – co znamienne – dokładnie tak, jak nakazuje dobra praktyka inżynierska: zaproponował, żeby najpierw odczytać z czujnika sam jego numer identyfikacyjny (każdy taki układ potrafi „przedstawić się” stałą liczbą) i potwierdzić,

że komunikacja w ogóle działa, a dopiero potem dopisywać resztę. To podejście „najpierw mały, działający dowód, potem reszta”. Był też bardziej rozmowny w tłumaczeniu decyzji i odnalazł świeże wątki z forum producenta sprzed kilku tygodni, których sam konstruktor nie znał.

Jeden komentarz uczestnika sesji dobrze streszcza różnicę: jeden agent świetnie nadaje się do tworzenia projektu od zera, drugi – do poprawiania i pracy na istniejącym kodzie. To nie spór „który jest lepszy”, lecz obserwacja, że różne narzędzia pasują do różnych faz pracy. I to ważne odkrycie w czasach, gdy same modele radzą sobie z coraz dłuższymi zadaniami: o wyniku coraz mniej decyduje „czy AI da radę”, a coraz bardziej „jak je poprowadzić”.

Konfabulacja na żywo. Najważniejszy moment eksperymentu nie był spektakularny – i właśnie dlatego jest pouczający. W pewnej chwili pierwszy agent oświadczył z pełnym przekonaniem, że konkretne wyprowadzenie mikrokontrolera (nóżka układu, do której podłączono diodę) jest ustawione jako wyjście. Tyle że człowiek pomylił się przy konfiguracji i ustawił je jako wejście. Zapytany wprost, skąd ma tę informację, agent przyznał, że „strzelił z pamięci na podstawie nazwy sygnału, zamiast sprawdzić w pliku konfiguracyjnym”. Gdyby po drugiej stronie siedział ktoś bez wiedzy o sprzęcie, błąd przeszedłby dalej – a potem zaczęłoby się wielogodzinne szukanie, czemu dioda nie świeci.

To konfabulacja w czystej postaci – i warto przy tym słowie się zatrzymać, bo będzie wracać. Konfabulacja (w żargonie AI częściej zwana „halucynacją”) to nie awaria ani komunikat o błędzie, ale wiarygodnie brzmiące, lecz fałszywe twierdzenie, podane z tą samą pewnością co prawda. Model nie „kłamie” w ludzkim sensie – generuje najbardziej prawdopodobny ciąg słów, a nie sprawdzony fakt. I tu dotykamy sedna zmiany w 2026 roku: narzędzia stały się na tyle dobre, że ich odpowiedzi brzmią coraz bardziej autorytatywnie – co paradoksalnie czyni konfabulacje groźniejszymi, bo trudniej je odruchowo podważyć.

Limit jako problem produkcyjny. Eksperyment skończył się w sposób, o którym milczą entuzjastyczne filmiki: nie na działającym programie wgranym do układu, lecz na wyczerpaniu limitu. Praca na najwydajniejszym modelu, w trybie największego wysiłku obliczeniowego, w abonamencie za około 20 dolarów miesięcznie wystarczyła na mniej więcej godzinę. Drugi agent zdążył jednak doprowadzić do tego, że temperatura i ciśnienie pojawiły się na ekranie i reagowały na dotknięcie czujnika – zadanie technicznie wykonane. Po drodze trzeba było m.in. dopisać program przeszukujący magistralę I²C, bo adres czujnika zależy od stanu jednej z jego nóżek, a agent początkowo przyjął zły. Wniosek konstruktora jest trzeźwy: realna praca komercyjna oznacza droższe abonamenty albo dokupowanie mocy obliczeniowej,

a uzależnienie procesu projektowego od narzędzia działającego „na porcję” to osobne, niedoceniane ryzyko.

Z tego przykładu wylaniają się cztery obserwacje, które nie zdezaktualizowały się wraz z postępem modeli, bo dotyczą nie samej zdolności AI, lecz sposobu pracy z nim. Po pierwsze, wiedza o sprzęcie pozostaje warunkiem koniecznym, a nie luksusem – to ona pozwoliła wychwycić zmyślenie. Po drugie, embedded jest pod względem dostępnego „materiału do nauki” w trudniejszym położeniu niż typowe programowanie. Po trzecie, różne narzędzia pasują do różnych faz pracy. Po czwarte, koszty i limity to realne ograniczenia, nie detal.

3. Co naprawdę zmieniło się w liczbach

Przedstawiony przykładowy eksperyment to jedna sesja jednej osoby. Żeby zobaczyć, czy jego wnioski są regułą, czy jednostkowym przypadkiem, trzeba je zestawić z danymi w skali, której pojedynczy konstruktor nigdy nie zbierze. A te dane w ciągu roku ułożyły się w obraz ciekawszy niż „było źle, jest dobrze”.

Adopcja niemal powszechna, zaufanie – nie. Najszerszy przekrój przez globalną społeczność programistów daje coroczna ankieta Stack Overflow; w edycji 2025 odpowiedziało w niej ponad 49 tysięcy osób ze 166 krajów. Z jednej strony korzystanie z AI stało się normą: 84% programistów używa lub planuje używać narzędzi AI, wobec 76% rok wcześniej [7]. Z drugiej – zaufanie do trafności tych narzędzi nie rośnie wraz z adopcją, lecz spada: odsetek ufających dokładności wyników obniżył się do 29%,

podczas gdy w 2024 roku wynosił 40%[8]. To nie sprzeczność, lecz dojrzewanie. Ludzie używają narzędzia coraz powszechniej, a zarazem coraz lepiej rozumieją, gdzie ono zawodzi – i przestają mu ufać „na słowo”.

Co istotne dla naszego tematu, najostrożniejsi są ci, którzy wiedzą najwięcej. To doświadczeni programiści wypadają w ankiecie najbardziej sceptycznie – najrzadziej deklarują wysokie zaufanie do wyników AI i najczęściej wprost im nie ufają [8]. Im większa wiedza i odpowiedzialność, tym mocniej wynik modelu traktuje się jako materiał do sprawdzenia, a nie gotowy produkt. To empiryczne potwierdzenie, że wiedza fachowa nie traci na znaczeniu, tylko zmienia rolę: z „autora każdej linijki” na „recenzenta i arbitra poprawności”.

Praca przesunęła się od pisania ku sprawdzaniu. Najlepszym dowodem, że rola się zmieniła, a nie zniknęła, jest to, na co programiści realnie poświęcają czas. Z badań rynkowych z początku 2026 roku wynika odwrócenie wcześniejszego układu: deweloperzy przeznaczają tygodniowo więcej godzin na przeglądanie kodu wygenerowanego przez AI niż na samodzielne pisanie nowego – około 11,4 wobec 9,8 godziny [9]. To konkretny, mierzalny ślad nowej roli: narzędzie produkuje, człowiek ocenia i koryguje. Te same badania pokazują, że entuzjastów czeka otrzeźwienie w dłuższym biegu – deklarowana produktywność rośnie o około jedną trzecią w pierwszych dwóch miesiącach, po czym wypłaszcza się około 180. dnia [9]. Pierwszy zachwyt jest realny, ale nie utrzymuje się w nieskończoność.

ADOPCJA KONTRA ZAUFANIE

Im powszechniejsze AI, tym mniej mu ufamy

Korzystanie rośnie

2024 **76%**
2025 **84%**

używa lub planuje używać narzędzi AI

Zaufanie spada

2024 **40%**
2025 **29%**

ufa dokładności wyników AI

Kto ufa najmniej?

Najbardziej sceptyczni są najbardziej doświadczeni programiści – to oni najrzadziej deklarują wysokie zaufanie do AI. Wiedza nie znika; zmienia rolę z autora w recenzenta.

Źródło: Stack Overflow Developer Survey 2025 (>49 tys. odpowiedzi, 166 krajów).

Elektronika Praktyczna

Rysunek 2. Adopcja kontra zaufanie. Korzystanie z AI stało się normą, ale zaufanie do trafności wyników spadło – najbardziej sceptyczni są najbardziej doświadczeni. Źródło: Stack Overflow Developer Survey 2025

„Odczucie kontra pomiar” – runda druga. W części 1. widzieliśmy, że w 2025 roku programiści czuli się szybsi, niż byli naprawdę. W 2026 ten wątek wrócił w dojrzałszej formie. W maju 2026 METR opublikował ankietę wśród specjalistów technicznych: deklarowali oni medianowo od 1,4 do 2-krotny wzrost wartości swojej pracy dzięki AI [10]. Liczba imponująca – ale autorzy zrobili coś rzadkiego: ostrzegli przed własnym wynikiem. To właśnie zespół METR podawał najniższe szacunki ze wszystkich badanych grup, prawdopodobnie dlatego, że jego pracownicy mają w pamięci wcześniejsze rozbieżności między postrzeganą a rzeczywistą produktywnością [10]. Morał jest ten sam co rok wcześniej, choć liczby urosły: deklaracje o „dwukrotnym przyspieszeniu” warto traktować jako odczucie, nie jako obiektywnie zmierzony fakt.

Co z tego wynika dla sporu z części 1. Zestawienie danych rozbraja oba dawne stanowiska, ale nie po równo. Teza „AI jest bezużyteczne” jest dziś nie do obrony – adopcja jest powszechna, narzędzia mierzalnie zdolniejsze, a praktycy zgodnie mówią, że dobrze prowadzony agent realnie „dowodzi”. Ale teza „AI zastąpiło programistów” rozбивa się o to, że najbardziej doświadczeni ufają mu najmniej, że ciężar pracy przesunął się ku sprawdzaniu, a nie zniknął, i że zysk pojedynczego człowieka nie przekłada się prosto na zysk zespołu. Spór nie został rozstrzygnięty na korzyść jednej ze stron – on dojrzał. A na nowe pytania świat sprzętu i embedded odpowiada inaczej niż świat aplikacji, którego dotyczy większość tych liczb.

4. „Embedded to nie web”: dlaczego konstruktor gra według innych reguł

Wszystkie liczby z poprzednich części opisują przede wszystkim świat typowego oprogramowania – aplikacji internetowych, usług sieciowych, kodu w popularnych językach jak Python czy JavaScript. To środowisko, w którym AI radzi sobie najlepiej, bo ma się na czym uczyć. Konstruktor elektronik pracuje w innym środowisku. I właśnie ta różnica, najczęściej pomijana w forumowych sporach, sprawia, że optymistyczne dane z 2026 roku nie przenoszą się wprost na jego biurko.

Skąd bierze się różnica: na czym model się uczył. Model językowy jest tak dobry, jak materiał, na którym go wytrenowano. Kodu aplikacji internetowych są w sieci miliardy linii – niezliczone przykłady, biblioteki, gotowe rozwiązania niemal każdego typowego problemu. Kodu sterującego konkretnym mikrokontrolerem czy konkretną wersją biblioteki producenta jest nieporównanie mniej; do tego jest rozproszony, często zamknięty w firmach jako tajemnica handlowa i szybciej się dezaktualizuje, bo co rok wychodzą nowe układy i wersje narzędzi. To nie brak, który pokona mądrzejsza wersja modelu – to strukturalna cecha

dziedziny. I dokładnie ona ujawniła się w eksperymencie z części 2.: gdy biblioteka była zbyt świeża, by trafić do materiału treningowego, model nie powiedział „nie wiem”, tylko zgadywał przez analogię do starszej wersji – i konfabulował.

Projektowanie układów cyfrowych: ten sam problem, ale zmierzony. Najlepiej udokumentowanym przypadkiem tej luki jest generowanie kodu, którym opisuje się układy cyfrowe – na przykład procesory czy układy programowalne FPGA. Służą do tego języki opisu sprzętu, w branży skracane do HDL (od ang. *hardware description language*); najpopularniejsze to Verilog i VHDL. To dziedzina jeszcze trudniejsza dla AI niż pisanie programów na mikrokontroler, ale właśnie dlatego najlepiej zbadać – a jej wnioski przenoszą się na całą „twardą” stronę elektroniki.

Literatura naukowa nazywa rzecz wprost: modele wykazują silniejszą skłonność do halucynacji przy językach opisu sprzętu niż przy językach ogólnego przeznaczenia, generując kod, który wygląda wiarygodnie, lecz zawiera błędy składniowe i funkcjonalne – głównie z powodu niedostatku danych specyficznych dla sprzętu w materiale treningowym [11]. To nie pojedyncza obserwacja, lecz uznany problem badawczy: redukcja tych zmyśleń doczekała się w 2025 roku własnych prac i dedykowanych metod. Jedną z nich, przedstawioną na czołowej konferencji branżowej, wprowadza systematyczną klasyfikację typów halucynacji w generowaniu Verilogu i mechanizm tłumaczenia opisów symbolicznych – takich jak tablice prawdy czy diagramy stanów – na precyzyjny język [12]. Sam fakt, że powstał osobny nurt badań nad tym, jak powstrzymać AI od zmyślania w kodzie sprzętowym, mówi więcej niż dowolna opinia z forum.

Rok 2026: przemysł odpowiada „weryfikacją w pętli”. Branża projektowania układów, zmuszona własnym problemem, wypracowała wzorzec wart przeniesienia do całej pracy z AI. Skoro modelowi nie można ufać „na słowo”, obudowuje się go narzędziami, które same sprawdzają jego wynik. Przykładem jest system, który firma Cadence wprowadziła w lutym 2026 roku – agentowy asystent najwcześniejszego, najbardziej ryzykownego etapu projektowania układu, gdzie zamiast konstruktora zamienia się w kod opisujący sprzęt. System automatyzuje tworzenie tego kodu, a także testbenchy – programów testujących, sprawdzających, czy projekt zachowuje się zgodnie z założeniami – oraz planów weryfikacji [13]. Kluczowy jest sposób, w jaki to robi: łączy klasyczną, rozwijaną od dekad analizę statyczną (sprawdzanie kodu narzędziami, które nie zgadują, lecz działają według ścisłych reguł) z rozumowaniem modelu językowego, dokładającego kontekst, jakiego ściśle narzędzia same nie mają [13]. To właśnie „weryfikacja w pętli” – AI nie jako wyrocznia, lecz jako

EMBEDDED TO NIE WEB

Dlaczego konstruktor gra według innych reguł

Typowe oprogramowanie

aplikacje webowe, Python, JavaScript

miliardy linii kodu

AI ma się na czym uczyć →
rzadziej zmyśla

Sprzęt i embedded

świeże biblioteki, HDL, konkretne układy

mało danych

AI zgaduje przez analogię →
częściej konfabuluje

Dowód: języki opisu sprzętu (HDL)

Modele halucynują w kodzie sprzętowym częściej niż w zwykłym — redukcja tych zmyśleń stała się w 2025 r. osobnym nurtem badań. Im bliżej sprzętu, tym wyższy „podatek na sprawdzanie”.

Źródło: prace nad ograniczaniem halucynacji w generowaniu Verilog/HDL (2025).

Elektronika Praktyczna

Rysunek 3. Embedded to nie web. Typowe oprogramowanie ma w sieci miliardy linii kodu do nauki; kodu sprzętowego jest mało. Stąd w embedded model częściej zgaduje i konfabuluje. Źródło: prace nad halucynacjami w HDL (2025)

jeden element procesu, w którym obok niego stoi sprawdzacz wyłapujący zmyślenia.

Ale „autonomiczny projektant” wciąż nie istnieje. Trzeba uczciwie ostudzić nadzieje, bo marketing 2026 roku lubi słowo „autonomiczny”. Wspólnym mianownikiem zapowiedzi wszystkich trzech największych producentów narzędzi projektowych jest model pracy „orkiestrowany przez AI, ale nadzorowany przez człowieka”, w którym rola inżyniera ewoluuje w stronę stratega definiującego cele, a nie wykonawcy każdego kroku [14]. A tam, gdzie zadanie robi się naprawdę trudne – przy projektowaniu najbardziej złożonych, przestrzennych struktur układów – szczerowość bywa większa: na branżowej konferencji w 2026 roku przyznano wprost, że pełny, autonomiczny agent do takiego projektowania i końcowej weryfikacji wciąż nie istnieje i nadal czeka na dopracowanie we współpracy z klientami [15]. To zdrowy kontrapunkt dla nagłówków o „końcu zawodu inżyniera”.

Co to znaczy w praktyce pracowni. Wniosek nie brzmi „AI jest w embedded bezużyteczne” – to ten sam błąd nadmiernego uogólnienia, który zarzucaliśmy obu obozom. Brzmi raczej: te same narzędzia, które w typowym programowaniu w 2026 roku bywają imponujące, w zadaniach sprzętowych schodzą o klasę niżej i wymagają znacznie gęstszej nadzoru. Im bliżej fizycznego układu – im bardziej liczą się dokładne czasy sygnałów, równoczesność zdarzeń i rejestry konkretnego układu scalonego – tym mniej można ufać wynikowi „na słowo”. Konstruktor jest

więc w sytuacji, w której potencjalny zysk z AI jest realny, ale „podatek” w postaci sprawdzania – wyższy niż u twórcy aplikacji internetowych. Dlatego optymistyczne liczby z poprzednich części powinniśmy czytać z poprawką: dotyczą głównie świata, w którym AI miało się na czym nauczyć. Jego świat dopiero do tego dorasta.

5. Poza kodem: reszta pracy w pracowni

Spór na forach toczy się prawie wyłącznie o programowanie, bo to najbardziej widowiskowe zastosowanie AI. Ale konstruktor spędza nad kodem tylko część dnia. Reszta to projektowanie płytek drukowanych, dokumentacja, przekopywanie się przez noty katalogowe, przygotowanie produkcji. I właśnie tutaj rok 2026 przyniósł najwięcej – z tą różnicą, że AI wchodzi w te obszary nie jako chatbot, do którego się pisze, lecz jako funkcja wbudowana w narzędzia, których konstruktor i tak używa.

Projektowanie płytek: AI już siedzi w narzędziach.

Trzeba zacząć od rozróżnienia, bo „projektowanie elektroniki” obejmuje dwa różne światy. Jeden to projektowanie układów scalonych. Drugi, bliższy większości czytelników, to projektowanie płytek drukowanych (PCB): rozmieszczenie elementów i poprowadzenie ścieżek łączących je na laminacie. Narzędzia do tego, zwane w branży EDA (od ang. *electronic design automation*, czyli komputerowe wspomaganie projektowania elektroniki), od kilku lat zyskują funkcje oparte na AI – a w 2026 roku tempo wyraźnie przyspieszyło.

Konkrety z tego roku robią wrażenie. Cadence, jeden z trzech największych dostawców narzędzi projektowych, rozbudował w marcu 2026 współpracę z firmą NVIDIA o – jak to ujęto – autonomiczne, długo działające agenty, które tłumaczą zamysł projektowy na zautomatyzowane procesy, generują projekty i usuwają błędy, zarządzając złożonymi przepływami pracy od początku do końca [16]. W warstwie samych płytek narzędzie tej firmy wnosi do projektowania PCB rozmieszczanie elementów i prowadzenie ścieżek sterowane przez AI [17], a deklarowane skrócenie czasu bywa liczone w wielokrotnościach. Takie liczby warto jednak czytać jako dane producenta, nie niezależny pomiar – z tą samą ostrożnością, z jaką w części 3. traktowaliśmy deklaracje o „dwukrotnej produktywności”.

Granice wyznaczają tu sami inżynierowie. Jak podsumowano w branżowym przeglądzie z 2026 roku, wspólnym kierunkiem wszystkich trzech wielkich dostawców jest przepływ pracy „orkiestrowany przez AI, ale nadzorowany przez człowieka”, w którym rola inżyniera ewoluuje w stronę stratega definiującego cele i zarządzającego systemem zajmującym się szczegółami wykonania [14]. To ta sama zmiana roli – z wykonawcy na recenzenta i decydenta – którą widzieliśmy u programistów. Charakter pracy się zmienia, ale odpowiedzialność za to, czy płytka da się wyprodukować, schłodzić i przetestować, zostaje po stronie człowieka.

Dokumentacja techniczna: najmniej efektowne, najbardziej praktyczne. Jeśli jest obszar, w którym AI daje konstruktorowi realny zysk przy stosunkowo niskim ryzyku, to praca z dokumentacją. Z jednej strony tworzenie: wstępne opisy działania, komentarze do kodu, szkice instrukcji, opisy interfejsów – jako materiał do redakcji przez człowieka, nie produkt końcowy. Z drugiej, jeszcze cenniejsza, analiza cudzej dokumentacji: zamiast ręcznie przewracać dwustustronicową notę katalogową w poszukiwaniu jednego parametru, można poprosić model o wyłuskanie konkretnej procedury albo zestawienie danych. Trzeba tu jednak postawić to samo ostrzeżenie co przy kodzie, bo mechanizm konfabulacji jest identyczny: model równie pewnie poda parametr, którego w nocie nie ma, jak ten, który tam jest. Dokumentacja tworzona lub streszczana przez AI wymaga sprawdzenia ze źródłem.

Produkcja i kontrola jakości. Na końcu łańcucha AI pracuje od dawna i najmniej kontrowersyjnie: w automatycznej optycznej inspekcji płytek (kamery z oprogramowaniem wychwytyjącym wady lutowania i montażu), w analizie defektów na obrazach, w optymalizacji linii. Rozpoznawanie obrazu to zadanie, w którym uczenie maszynowe jest sprawdzone od lat – i dobrze pokazuje, że „AI w elektronice” to nie tylko modni agenci kodowania, ale i dojrzałe, wąsko specjalistyczne narzędzia robiące jedną rzecz porządnie.

Edge AI: konstruktor po drugiej stronie. Jest wreszcie obszar, w którym AI nie jest narzędziem konstruktora, lecz przedmiotem jego pracy: modele uruchamiane lokalnie, bezpośrednio w urządzeniu – na mikrokontrolerze, układzie SoC czy FPGA – zamiast w chmurze. To podejście, zwane edge AI (przetwarzanie „na brzegu” sieci, tam gdzie powstają dane, a nie na odległym serwerze), jest dziś silnym trendem. Konstruktor musi w nim połączyć klasyczne kompetencje sprzętowe ze zrozumieniem, jak model działa, ile potrzebuje pamięci i mocy oraz jak go zmieścić w ograniczonym układzie. Koło się domyka, bo część narzędzi projektowych zaczyna wspierać projektowanie płytek właśnie pod kątem wbudowanych modeli AI.

Wspólny mianownik tej części jest taki sam jak przy kodzie. AI najlepiej radzi sobie z zadaniami dobrze zdefiniowanymi i weryfikowalnymi – optymalizacją rozmieszczenia pod jednoznaczne kryteria, rozpoznawaniem wad na obrazie, wyszukiwaniem w dokumentacji. Najślabiej tam, gdzie trzeba twórczej decyzji obejmującej cały system albo wiedzy, której w danych nie było. I wszędzie wartość pojawia się dopiero wtedy, gdy po drugiej stronie stoi inżynier potrafiący ocenić wynik.

6. Ryzyka i granice: tu spór dojrzał najbardziej

Jeśli w jednym miejscu widać, że rok 2026 zmienił dyskusję o AI, to właśnie tutaj. W 2025 roku ryzyka były w dużej mierze przewidywaniami; dziś są policzone. To najważniejsza zmiana w całym sporze: postęp w zdolnościach modeli nie tylko nie usunął zagrożeń, ale część z nich pogłębił – bo im sprawniejszy i bardziej przekonujący asystent, tym łatwiej przeoczyć moment, w którym się myli.

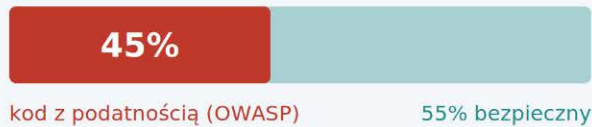
Konfabulacja to cecha, nie usterka. Zmyślanie nie jest błędem, który da się „naprawić” lepszym poleceniem czy nowszą wersją modelu. Wynika z samej zasady działania: model generuje najbardziej prawdopodobny ciąg słów pasujący do kontekstu, a nie zweryfikowaną prawdę – dlatego równie płynnie poda parametr, który istnieje, jak ten, którego nigdy nie było. W embedded, jak pokazała część 4., skłonność do konfabulacji jest dodatkowo większa. Co gorsza, postęp 2026 roku działa na niekorzyść czujności: im lepszy model, tym rzadziej się myli – ale tym bardziej autorytatywnie brzmi, gdy już zmyśli. Wniosek: każdy wynik dotyczący sprzętu – rejestr, czas sygnału, adres, parametr z noty – traktuje się jako twierdzenie do sprawdzenia, nie jako fakt.

Bezpieczeństwo kodu: cena prędkości. Tu dane z 2026 roku są najbardziej wymowne i najmocniej studzą huraoptymizm. Przypomnijmy liczbę z części 1. i dołożmy resztę obrazu: badanie firmy Veracode na ponad stu modelach wykazało, że 45% kodu generowanego przez AI zawierało znaną lukę bezpieczeństwa, a odsetek rozwiązań wolnych od takich błędów utknął na poziomie około

CENA PRĘDKOŚCI

Kod AI działa coraz lepiej — ale nie bezpieczniej

Co drugi fragment z luką



odsetek bezpiecznego kodu utknął na ~55% mimo poprawy samego działania (2025–2026)

Kod współtworzony przez AI

1,6×	więcej znalezisk bezpieczeństwa
1,75×	więcej błędów logicznych
2,74×	częściej podatność typu XSS
+107%	podatności na bazę kodu (rok do roku)

Dlaczego to ważne

Bezpieczeństwa nie widać w „działa / nie działa” — trzeba je celowo sprawdzić.
W urzędzeniu pracującym latami to dług, który ktoś spłaci — często już po wdrożeniu.

Rysunek 4. Cena prędkości. Mniej więcej co drugi fragment kodu AI zawiera znaną lukę, a odsetek bezpiecznego kodu nie poprawił się mimo lepszego działania. Kod współtworzony przez AI ma też więcej błędów. Źródła: Veracode, CodeRabbit, OSSRA (2025–2026)

55% przez cały okres 2025–2026 — mimo wyraźnej poprawy samego działania [6]. Sedno: modele nauczyły się pisać kod, który działa, ale nie nauczyły się pisać go bezpiecznie — bo bezpieczeństwa nie widać w „działa / nie działa”, trzeba je celowo sprawdzić. Niezależne analizy potwierdzają kierunek z różnych stron: kod współtworzony przez AI dawał około 1,6 raza więcej znalezisk bezpieczeństwa i 1,75 raza więcej błędów logicznych niż kod ludzki, a podatność umożliwiającą wstrzyknięcie obcego skryptu zawierał blisko trzykrotnie częściej [18]. W skali całych projektów widać to jako lawinę: raport badający otwarte komponenty oprogramowania odnotował wzrost liczby podatności na bazę kodu o 107% rok do roku [19]. Dla konstruktora, którego kod trafia do urzędzenia pracującego latami i trudnego do aktualizacji, to realny dług, który ktoś kiedyś będzie musiał spłaci — często już po wdrożeniu.

„Śmiercionośna triada”: ryzyko rosnące wraz z autonomią. Im bardziej samodzielni stają się agenci — a w 2026 roku to główny kierunek — tym poważniejsze staje się zagrożenie, które łatwo przeoczyć, bo nie przypomina klasycznego włamania. Simon Willison nazwał je „śmiercionośną triadą”: niebezpieczną sytuacją, w której narzędzie AI ma jednocześnie trzy zdolności — dostęp do prywatnych danych (może czytać twoje pliki, dokumentację, repozytoria), kontakt z treścią z niezauważanego źródła (przetwarza coś, co przyszło z zewnątrz — stronę internetową, e-mail, cudzy dokument) oraz możliwość wysłania danych na zewnątrz. Jeśli system ma wszystkie trzy naraz, jest podatny na atak [20]. Mechanizm jest podstępny, bo nie wymaga dziury w klasycznym kodzie: w treści z zewnątrz

można ukryć polecenie (tzw. prompt injection — przemycenie do modelu instrukcji udającej część danych), które agent posłusznie wykona, np. wysyłając poufne dane projektowe w nieznane miejsce [21]. To nie teoria: w 2025 roku odnotowano realne ataki tego typu na firmowe systemy. Nie bez powodu w ankietach z początku 2026 roku przemycanie poleceń awansowało na drugie miejsce wśród najbardziej dotkliwych problemów pracy z agentami [9]. Dla działu R&D płyną stąd dwa wnioski: im więcej uprawnień damy agentowi, tym ostrożniej trzeba projektować, co wolno mu zrobić — a tam, gdzie w grę wchodzi tajemnica handlowa, poważnie rozważyć model działający lokalnie.

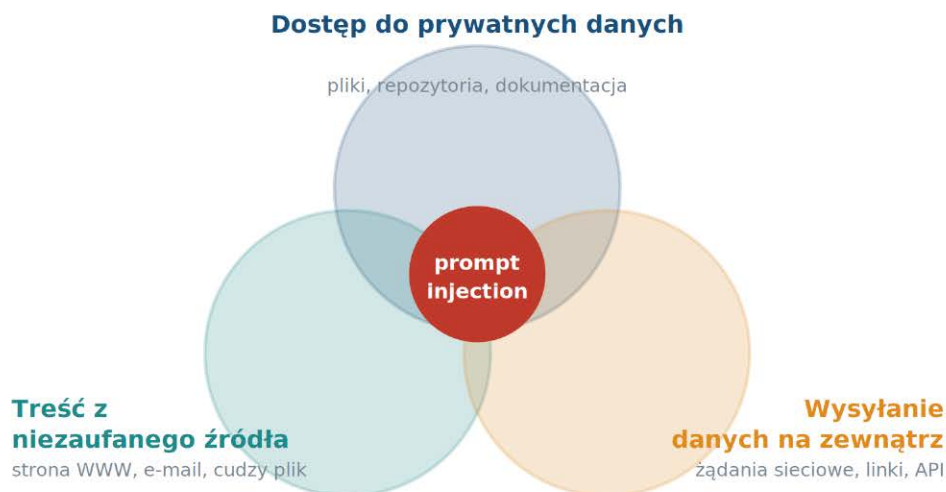
Erozja kompetencji. Ryzyko subtelniejsze, dotyczące ludzi, nie kodu — i w 2026 roku, gdy AI jest powszechne już na etapie nauki, ważniejsze niż kiedykolwiek. Badania nad nauką programowania z pomocą AI pokazują powtarzalny wzorec: uczący się kończą zadania szybciej, ale ich rozwiązania są bardziej do siebie podobne i płytsze pod względem zrozumienia. Dla doświadczonego konstruktora to mniejszy problem; dla wchodzących do zawodu — realne zagrożenie, bo zdolność wychwycenia konfabulacji bierze się właśnie z tej wiedzy, którą AI pozwala obejść. Tu kryje się odpowiedź na forumowe „po co kursy, skoro AI wszystko pisze”: bez własnego fundamentu nie sposób ocenić, czy AI pomaga, czy prowadzi w ślepą uliczkę. Im lepsze narzędzia, tym umiejętność oceny — a nie samego pisanie — staje się prawdziwym zawodem.

Koszty i uzależnienie od narzędzia. Ryzyko prozaiczne, które ujawnił eksperyment z części 2., a które dane z 2026 roku potwierdzają jako problem numer jeden. W ankietach z początku roku zmienność

ŚMIERCIONOŚNA TRIADA

Kiedy samodzielny agent staje się groźny

Ryzyko pojawia się, gdy narzędzie AI ma jednocześnie wszystkie trzy zdolności naraz:



Mając wszystkie trzy naraz, agent może wykonać polecenie ukryte w treści z zewnątrz — np. wysłać poufne dane projektowe. Bez żadnej dziury w klasycznym kodzie.

Rada: ogranicz uprawnienia agenta; przy tajemnicy handlowej — model lokalny.

Źródło: S. Willison, „The lethal trifecta for AI agents” (2025).

Elektronika Praktyczna

Rysunek 5. Śmiercionośna triada. Gdy agent ma naraz dostęp do prywatnych danych, kontakt z treścią z zewnątrz i możliwość wysyłania danych, staje się podatny na przemycone polecenia – bez żadnej dziury w klasycznym kodzie. Źródło: S. Willison (2025)

kosztów wysunęła się na pierwsze miejsce wśród bolączek pracy z agentami – przy rozliczaniu za zużycie miesięczne rachunki potrafią wahać się dwu- lub trzykrotnie z kwartału na kwartał [9]. Do tego dochodzi uzależnienie procesu projektowego od zewnętrznej usługi, jej cennika, dostępności i polityki. Dla małego zespołu czy jednoosobowej pracowni to czynnik, który trzeba wkalkulować, zanim wpięcie AI w stały tryb pracy stanie się nieodwracalne.

Suma tych ryzyk nie prowadzi do wniosku „lepiej nie tykać” – to anachronizm w świecie, w którym z AI korzysta dziewięciu na dziesięciu programistów. Prowadzi do wniosku, że im zdolniejsze narzędzie, tym ważniejsze są ramy jego użycia: świadomość, gdzie model konfabuluje, jak go zabezpieczyć, jak pilnować jakości, kosztów i własnych kompetencji.

7. Jak korzystać rozsądnie: praktyki, które działają

Skoro w połowie 2026 roku spór o to, czy AI potrafi, jest w dużej mierze rozstrzygnięty, ciężar problemu przenosi się na to, jak z niego korzystać. Praktyki, które oddzielają pożytek od szkody, są dziś dobrze rozpoznane – wynikają wprost z danych i relacji przytoczonych wcześniej. I większość z nich nie zestarzała się wraz z postępem

modeli, bo dotyczy sposobu pracy, a nie konkretnej wersji narzędzia.

Świadomy tryb pracy zamiast „pisania na czuja”. Punktem wyjścia jest rozróżnienie dwóch trybów pracy z AI. Pierwszy bywa nazywany „pisananiem na czuja” (ang. *vibe coding*) – generowanie kodu, którego się nie czyta i nie sprawdza, byle działał. To dopuszczalne wyłącznie tam, gdzie błąd nic nie kosztuje: prototyp na jeden wieczór, sprawdzenie pomysłu, jednorazowy skrypt. Wszystko, co trafia do urządzenia albo produktu, wymaga trybu drugiego – inżynierskiego: narzędzie proponuje, ale człowiek odpowiada za architekturę, jakość i poprawność i każdy fragment przepuszcza przez własną ocenę. Praktyczna pętla wygląda jak zachowanie metodycznego agenta z części 2.: najpierw plan i mały, działający dowód, potem rozbudowa małymi krokami, na końcu testy i przegląd przez człowieka. To nie spowolnienie dla zasady – to ten sam „podatek na sprawdzanie”, którego pominięcie i tak wraca w postaci godzin spędzonych na przeglądach.

Mapa: gdzie AI realnie pomaga, a gdzie zawodzi. AI jest dziś mocne w zadaniach szablonowych i dobrze określonych: tworzeniu powtarzalnego „rusztowania” kodu, szkieletów sterowników, skryptów pomocniczych, programów testujących, pierwszych wersji dokumentacji,

wyszukiwaniu w długich notach, poznawaniu nieznanego interfejsu, optymalizacji rozmieszczenia pod jednoznaczne kryteria. AI jest słabe i wymaga gęstego nadzoru tam, gdzie brakuje materiału do nauki (świeże układy peryferyjne, nowe wersje bibliotek, niszowe rozwiązania), gdzie liczą się dokładne czasy sygnałów i równoczesność zdarzeń, gdzie problem wymaga twórczej decyzji obejmującej cały system, oraz wszędzie tam, gdzie błąd jest krytyczny dla bezpieczeństwa. Reguła kciuka pozostaje ta sama mimo postępu modeli: im bliżej fizycznego sprzętu i im świeższy ekosystem, tym mniej ufać wynikowi „na słowo”.

Jak wykrywać konfabulacje. Pytaj model wprost o źródło twierdzenia – „z którego rejestru to odczytałeś?”, „gdzie w nocy jest ten parametr?”; przyznanie się do „strzału z pamięci” zdarza się częściej, niż można by sądzić. Każdy parametr sprzętowy konfrontuj z dokumentacją źródłową, nie z odpowiedzią modelu. Traktuj to, że kod się kompiluje, jako warunek konieczny, ale niewystarczający – bo, jak pokazały dane o bezpieczeństwie, kod może działać i wciąż zawierać lukę. I nie oddawaj AI ostatniego kroku, czyli fizycznego uruchomienia na sprzęcie: to moment, w którym zmyślenie spotyka się z rzeczywistością i w którym potrzebny jest człowiek rozumiejący, co się dzieje.

Dawaj kontekst projektu. Wiele błędów bierze się stąd, że model zgaduje, bo nie dostał kontekstu. Warto mu go dostarczyć: plik z konwencjami projektu (nazewnictwo,

struktura, używana wersja biblioteki), jasne wskazanie wariantu układu, formatu danych, ograniczeń. Metodyczne narzędzie z eksperymentu samo o to dopytywało – i właśnie to odróżniało udaną sesję od konfabulacji. Im mniej model musi zgadywać, tym rzadziej zmyśla.

Zabezpiecz agenta i dane. Jeśli narzędzie ma jednocześnie dostęp do repozytorium, dokumentacji i sieci, pamiętaj o „śmiercionośnej triadzie”: ograniczaj uprawnienia, zachowaj szczególną ostrożność przy operacjach zmieniających coś na podstawie treści z zewnątrz, a tam, gdzie w grę wchodzi tajemnica handlowa, rozważ model działający lokalnie. Bezpieczeństwo agenta sprowadza się w praktyce do pilnowania, co w ogóle wolno mu zrobić.

Panuj nad kosztami. Skoro zmienność rachunków stała się w 2026 roku bólem numer jeden, warto z góry ustalić limity, obserwować zużycie i nie uzależniać krytycznych etapów projektu od jednego, rozliczanego za zużycie narzędzia. Zdrowa zasada: traktować AI jak płatne źródło mocy, a nie darmową wodę z kranu.

Te praktyki nie wymagają specjalnych narzędzi – wymagają nawyku traktowania AI jak bardzo szybkiego, ale niefrasobliwego współpracownika, którego pracę zawsze się sprawdza. Konstruktor, który tak pracuje, korzysta z realnego – i z miesiąca na miesiąc większego – przyspieszenia, nie płacąc za nie jakością, bezpieczeństwem ani własnymi kompetencjami.

MAPA ZASTOSOWAŃ

Gdzie AI realnie pomaga, a gdzie zawodzi

✓ AI pomaga	✗ AI zawodzi / pod nadzorem
• powtarzalne „rusztowanie” kodu • szkielety sterowników, skrypty • programy testujące (testbenche) • pierwsze wersje dokumentacji • wyszukiwanie w długich notach • poznawanie nieznanego interfejsu • optymalizacja pod jasne kryteria	• świeże układy bez danych • nowe wersje bibliotek • ścisły timing, równoczesność • decyzje obejmujące cały system • cokolwiek krytycznego dla bezpieczeństwa • rejestry konkretnej kości • fizyczne uruchomienie na sprzęcie

Reguła kciuka: im bliżej fizycznego sprzętu i im świeższy ekosystem, tym mniej ufać „na słowo”.

Oprac. na podstawie danych i praktyk omówionych w artykule (2025–2026).

Elektronika Praktyczna

Rysunek 6. Mapa zastosowań. Praktyczne zestawienie zadań, w których AI realnie pomaga, i tych, gdzie zawodzi lub wymaga gęstego nadzoru – przydatne nad biurkiem konstruktora

Zakończenie

Wróćmy do komentarza spod reklamy kursu: „po co kursy, skoro AI wszystko pisze”. Po przejściu przez dane z 2025 i 2026 roku widać, że pytanie jest źle postawione – i że odpowiedź na nie zmieniała się w ciągu zaledwie roku w sposób, który powinien uspokoić jedną stronę sporu i otężyć drugą. Tak, AI pisze coraz więcej i coraz lepiej; zarzut „to bezużyteczna zabawka” jest dziś nie do obrony. Ale nie, nie zastąpiło to konstruktora – przesunęło tylko ciężar jego pracy z pisania ku ocenianiu, a wartość jego wiedzy z „umiem to napisać” ku „umiem rozpoznać, kiedy to jest napisane źle”. Kursy i praktyka nie tracą sensu; przeciwnie, stają się tym, co odróżnia użytkownika AI od jej ofiary.

Prawdziwy obraz na połowę 2026 roku jest taki: spór nie został wygrany, on dojrzał. Pytanie „czy AI potrafi” ustąpiło pytaniom trudniejszym i ważniejszym – jakim kosztem dla bezpieczeństwa, jakości w długim horyzoncie i własnego zrozumienia; gdzie dokładnie w łańcuchu projektowym to się opłaca. Dla konstruktora elektronika odpowiedź jest przy tym wyraźnie inna niż dla twórcy aplikacji internetowych: jego świat – świeże układy, dokładne czasy

sygnałów, fizyczny sprzęt – jest dla AI trudniejszy, bo modele nie miały się na czym uczyć, więc „podatek na sprawdzanie” pozostaje u niego wyższy. To nie powód, by AI nie używać. To powód, by używać go świadomie.

Co do kierunku: skoro w 2026 roku narzędzia stają się coraz bardziej samodzielne, a zarazem rosną obawy o koszty, bezpieczeństwo i tajemnicę handlową, naturalnym następnym krokiem wydaje się przeniesienie ich bliżej konstruktora – z chmury na lokalne serwery i stacje robocze, podobnie jak modele AI przeszły z chmury wprost do urządzeń. Lokalny „agent inżynierski”, pracujący na danych firmy bez wypuszczania ich na zewnątrz, odpowiadałby na trzy bóle naraz: koszt, bezpieczeństwo i poufność. Czy tak się stanie i jak szybko – zobaczymy; w dziedzinie, w której w ciągu roku „raz działa, raz nie” zmieniło się w „działa”, ostrożność w prognozach jest najmądrzejszą postawą.

Jedno wydaje się pewne. Praca konstruktora elektronika się zmienia – z pisania każdej linijki i rysowania każdej ścieżki w stronę projektowania, oceniania i korygowania tego, co proponuje narzędzie.

WM, JS

Bibliografia i źródła

- Większość pozycji pochodzi z lat 2025–2026; przy danych producentów narzędzi (skrótowa funkcja) wskazana ostrożność – to deklaracje dostawcy, nie niezależny pomiar.
- METR, „Measuring the Impact of Early-2025 AI on Experienced Open-Source Developer Productivity” (RCT, lipiec 2025) – doświadczeni deweloperzy o 19% wolniejsi z AI, choć szacowali 20% przyspieszenia. <https://metr.org/blog/2025-07-10-early-2025-ai-experienced-os-dev-study/>
- METR, „We are Changing our Developer Productivity Experiment Design” (luty 2026) – próba powtórzenia badania nieukończona, bo deweloperzy odmawiali pracy bez AI; prawdopodobne przyspieszenie na początku 2026. <https://metr.org/blog/2026-02-24-uplift-update/>
- Omówienie danych METR (Horizon / Kwa i in.) i paradoksu wydajności: horyzont zadań z -60 min do -719 min (Claude Opus 4.6, luty 2026), podwojenie -128 dni; >90% adopcji, -27% kodu produkcyjnego, wzrost wydajności organizacji do -10%. <https://philipdubach.com/posts/93-of-developers-use-ai-coding-tools-productivity-hasnt-moved/>
- S. Willison, „An AI state of the union” (wywiad, Lenny’s Newsletter, 2026) oraz blog SimonWillison.net – listopad 2025 jako próg przejścia agentów z „w większości działa” do „działa”. <https://www.lennysnewsletter.com/p/an-ai-state-of-the-union>
- Veracode, „GenAI Code Security Report” (2025–2026), omówienie: 45% kodu generowanego przez AI zawiera podatność z listy OWASP; odsetek bezpiecznego kodu stabilny przy -55% mimo poprawy poprawności funkcjonalnej. <https://www.softwareseni.com/why-45-percent-of-ai-generated-code-contains-security-vulnerabilities/>
- Analiza utrzymywania się -55% bezpiecznego kodu w cyklu 2025–2026 mimo rosnącej poprawności funkcjonalnej (Veracode, 150+ modeli). <https://www.softwareseni.com/91-5-percent-of-vibe-coded-apps-have-vulnerabilities-and-what-the-q1-2026-research-actually-shows/>
- Stack Overflow, „2025 Developer Survey” (komunikat, 29.07.2025) – 84% adopcji (z 76% w 2024); 49 tys. odpowiedzi ze 166 krajów. <https://stackoverflow.com/company/press/archive/stack-overflow-2025-developer-survey/>
- Stack Overflow, „2025 Developer Survey – AI” (rozszerzona analiza, 29.12.2025) – zaufanie do trafności AI spadło do 29% (z 40%); doświadczeni deweloperzy najbardziej sceptyczni. <https://stackoverflow.blog/2025/12/29/developers-remain-willing-but-reluctant-to-use-ai-the-2025-developer-survey-results-are-here/>
- Badanie adopcji narzędzi AI Q1 2026 – więcej godzin na przeglądanie kodu AI (11,4) niż na pisanie nowego (9,8); plateau produktywności po -180 dniach; zmienność kosztów (#1) i prompt injection (#2) jako główne bolączki. <https://www.digitalapplied.com/blog/ai-coding-tool-adoption-2026-developer-survey>
- METR, „Measuring the Self-Reported Impact of Early-2026 AI on Technical Worker Productivity” (maj 2026) – deklarowane 1,4–2× wartości pracy; zespół METR podaje najniższe szacunki ze wszystkich grup. <https://metr.org/blog/2026-05-11-ai-usage-survey/>
- HDLCoRe: „A Training-Free Framework for Mitigating Hallucinations in LLM-Generated HDL” (arXiv 2503.16528, 2025) – silniejsza skłonność do halucynacji w HDL z powodu niedoboru danych; wątek ochrony własności intelektualnej. <https://arxiv.org/abs/2503.16528>
- Y. Yang i in., „HaVen: Hallucination-Mitigated LLM for Verilog Code Generation Aligned with HDL Engineers” (DATE 2025, arXiv 2501.04908) – taksonomia halucynacji i tłumaczenie reprezentacji symbolicznych. <https://arxiv.org/abs/2501.04908>
- Cadence, „ChipStack AI Super Agent” (luty 2026) – agentowy system do front-endu projektowania: generowanie RTL, testbenchy i planów weryfikacji, łączenie analizy statycznej z rozumowaniem LLM. <https://www.hpcwire.com/2026/02/12/cadence-introduces-agentic-ai-system-for-chip-design-and-verification/>
- „A Look at Agentic AI in the EDA Engineering Workflow”, Embedded (kwiecień 2026) – wspólny model „orkiestrowany przez AI, nadzorowany przez człowieka”; inżynier jako strateg definiujący cele. <https://www.embedded.com/a-look-at-agentic-ai-in-the-eda-engineering-workflow/>
- „CadenceLIVE 2026 – Can Agentic AI Finally Crack 3D IC Design Automation?”, Futurum (kwiecień 2026) – pełny agent do projektowania 3D IC i wielofizycznego signoff wciąż nie istnieje, czeka na współpracę z klientami. <https://futurumgroup.com/insights/cadencelive-2026-can-agentic-ai-finally-crack-3d-ic-design-automation/>
- Cadence–NVIDIA, „Accelerated Engineering Solutions for Agentic AI Chip and System Design” (marzec 2026) – autonomiczne, długo działające agenty w procesie projektowym. <https://www.hpcwire.com/off-the-wire/cadence-and-nvidia-unveil-accelerated-engineering-solutions-for-agentic-ai-chip-and-system-design/>
- Cadence, „Agentic AI for Chip Design / AI for Design” – rozmieszczanie i routing PCB sterowane przez AI (Allegro X AI). https://www.cadence.com/en_US/home/ai/ai-for-design.html
- CodeRabbit, „State of AI vs Human Code Generation Report” (grudzień 2025) – kod współtworzony przez AI: -1,57× więcej znalezisk bezpieczeństwa, 1,75× więcej błędów logicznych, 2,74× częściej podatność typu XSS. <https://www.kusari.dev/blog/ai-coding-assistants-in-2026-4x-faster-10x-riskier-the-hidden-security-cost>
- Black Duck, „OSSRA 2026” (omówienie) – wzrost liczby podatności na bazę kodu o 107% rok do roku. <https://philipdubach.com/posts/93-of-developers-use-ai-coding-tools-productivity-hasnt-moved/>
- S. Willison, „The lethal trifecta for AI agents: private data, untrusted content, and external communication” (16.06.2025). <https://simonwillison.net/2025/Jun/16/the-lethal-trifecta/>
- S. Willison, „New prompt injection papers: Agents Rule of Two and The Attacker Moves Second” (2.11.2025) – mechanizm prompt injection i rozszerzenie triady o zmianę stanu. <https://simonwillison.net/2025/Nov/2/new-prompt-injection-papers/>



Zaprojektowane z myślą o efektywności energetycznej

Dostrzeż każdy wat, oszczędzaj każdy mikroamper

Masz dość zużywania energii zasilającej tylko po to, by móc ją mierzyć? Chcesz otrzymywać na bieżąco alerty, aby chronić system bez wpływu na wydajność? Dzięki PAC1711 i PAC1811 monitorowanie zużycia energii przez system nie wpływa na czas pracy na baterii.

Układy PAC1711 i PAC1811 zmniejszają zużycie energii o 50% w porównaniu z typowymi 12-bitowymi i 16-bitowymi cyfrowymi monitorami. Pozwalają monitorować stan zasilania bez rozładowywania baterii, maksymalizując czas pracy i wydajność w każdej aplikacji zależnej od baterii lub zasilania awaryjnego.

Najważniejsze cechy

- Ultraniskie zużycie energii: PAC1711 pobiera zaledwie 100 μA , czyli zaledwie 50% prądu w podobnych monitorach 12-bitowych, a PAC1811 jedynie 140 μA , czyli zaledwie 50% prądu w podobnych monitorach 16-bitowych.
- Alerty systemowe w czasie rzeczywistym: Do 2 wyjść alarmowych do natychmiastowego powiadomienia o przekroczeniu limitu mocy.
- Łatwa integracja: PAC1711 i PAC1811 są kompatybilne pod względem pinów i obudowy z obudową SOT23-8, co ułatwia modernizację i migrację.
- Precyzyjne monitorowanie bez utraty rozdzielczości przy wyższych częstotliwościach próbkowania zarówno w PAC1711 (12 bitów), jak i PAC1811 (16 bitów).
- Lepsza ochrona systemu: Natychmiastowe alerty pomagają chronić urządzenia.

Ulepsz swój kolejny projekt o PAC1711 i PAC1811, gdzie wydajność, bezpieczeństwo i łatwość integracji łączą się w jedno.



microchip.com/MorePowerInsight

eprasa.pl 1761d60f5c

Nazwa i logo Microchip oraz logo Microchip są zastrzeżonymi znakami towarowymi Microchip Technology Incorporated w Stanach Zjednoczonych i innych krajach. Wszystkie pozostałe znaki towarowe stanowią własność ich zarejestrowanych właścicieli. © 2026 Microchip Technology Inc. Wszelkie prawa zastrzeżone. MEC2636A-POL-05-26

Zephyr RTOS – pierwsze 48 godzin: od Blinky do BLE

Praktyczny przewodnik po drodze, którą każdy projektant przechodzi tylko raz: od pustego terminala, przez migającą diodę, aż po działające peryferium Bluetooth Low Energy.

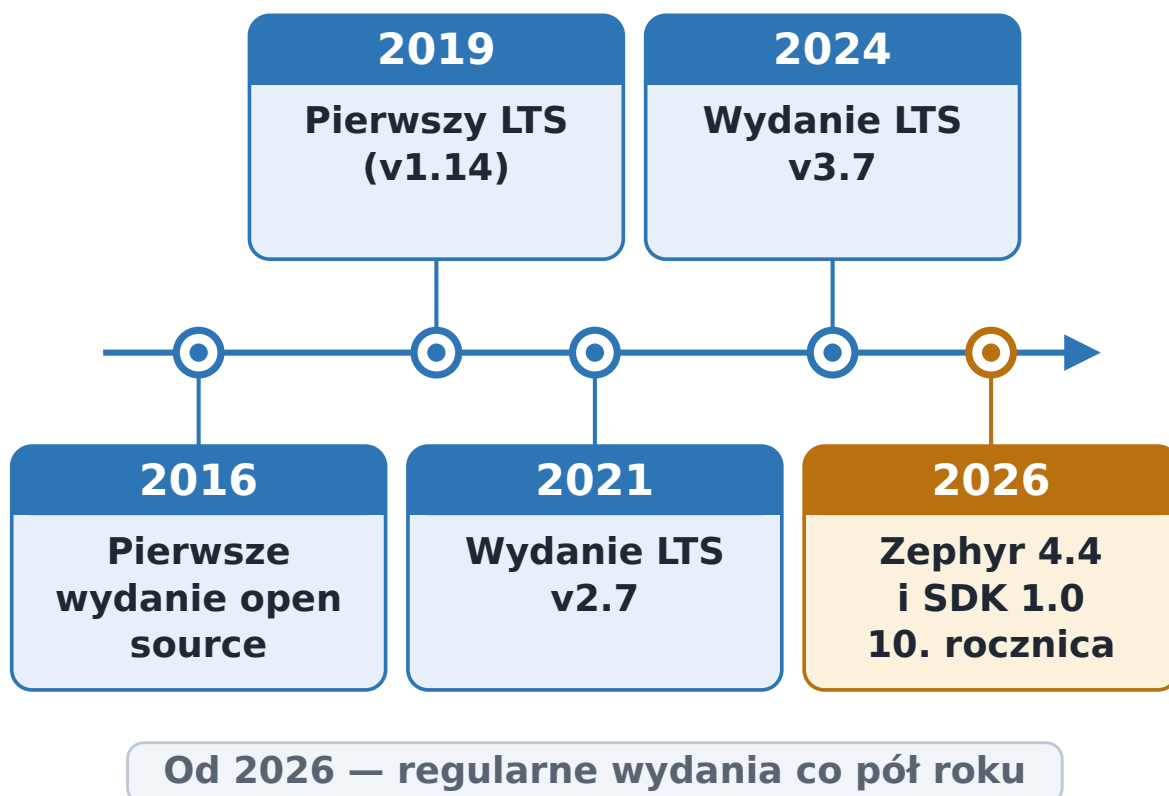
Zephyr RTOS skończył właśnie dziesięć lat i z niszowego projektu stał się jednym z filarów współczesnego oprogramowania wbudowanego. Nordic Semiconductor zbudował na nim swój nRF Connect SDK, a po jego kod sięgają między innymi Google, Meta, Intel i NXP. Dla projektanta przyzwyczajonego do programowania bare-metal albo do FreeRTOS pierwsze spotkanie z Zephyrem bywa jednak zaskakująco trudne – nie dlatego, że system jest źle zaprojektowany, lecz dlatego, że wymaga zmiany sposobu myślenia o tym, czym w ogóle jest projekt na mikrokontroler.

Dlatego potraktowaliśmy ten materiał jak krótki kurs wprowadzający – do przerobienia na płytce, nie tylko do przeczytania. Prowadzi on przez pierwsze 48 godzin pracy z Zephyrem: od pustego terminala, przez migającą diodę, aż po działające peryferium Bluetooth Low Energy. Kurs zaczyna się od kompletnej wyprawki kursanta (rozdział 3), prowadzi przez gotowe do powtórzenia listingi i kończy zestawem wskazówek praktyków. To nie jest wyczerpujący wykład o Zephyrze, lecz tylko przewodnik po pierwszym – i najtrudniejszym – odcinku drogi; przejście przez niego decyduje, czy zostaniemy z tym systemem na dłużej. Za przewodników służą nam czterej eksperci, którzy ten odcinek przeszli przed nami i nauczyli go tysiące innych: Kate Stewart, Shawn Hymel, Carles Cufí i Mohammad Afaneh.

1. Dlaczego akurat Zephyr – i dlaczego „48 godzin”

Zephyr to otwarty system operacyjny czasu rzeczywistego (RTOS – *Real Time Operating System*) rozwijany pod auspicjami Linux Foundation. Jego pierwsze wydanie ukazało się w lipcu 2016 roku, więc w 2026 projekt obchodzi okrągłą, dziesiątą rocznicę (rysunek 1).

To już nie jest eksperyment. Kate Stewart – dyrektorka projektu Zephyr i wiceprezes Linux Foundation odpowiedzialna za niezawodne systemy wbudowane – przy okazji jubileuszu ujęła pierwotny cel projektu wprost: chodziło o to, by zespoły mogły budować niezawodne systemy czasu rzeczywistego bez uzależnienia od jednego dostawcy, jednego łańcucha narzędzi i jednego zamkniętego stosu



Rysunek 1. Dekada Zephyra (2016–2026). Projekt wydaje główne wersje co pół roku; wydania LTS – utrzymywane długoterminowo – zaznaczono osobno. Stan na wydanie 4.4 z kwietnia 2026 roku

oprogramowania. Po dekadzie, jak dodaje Stewart, Zephyr jest dojrzałym i godnym zaufania RTOS-em, wbudowanym w produkty projektowane z myślą o latach, a często dekadach eksploatacji.

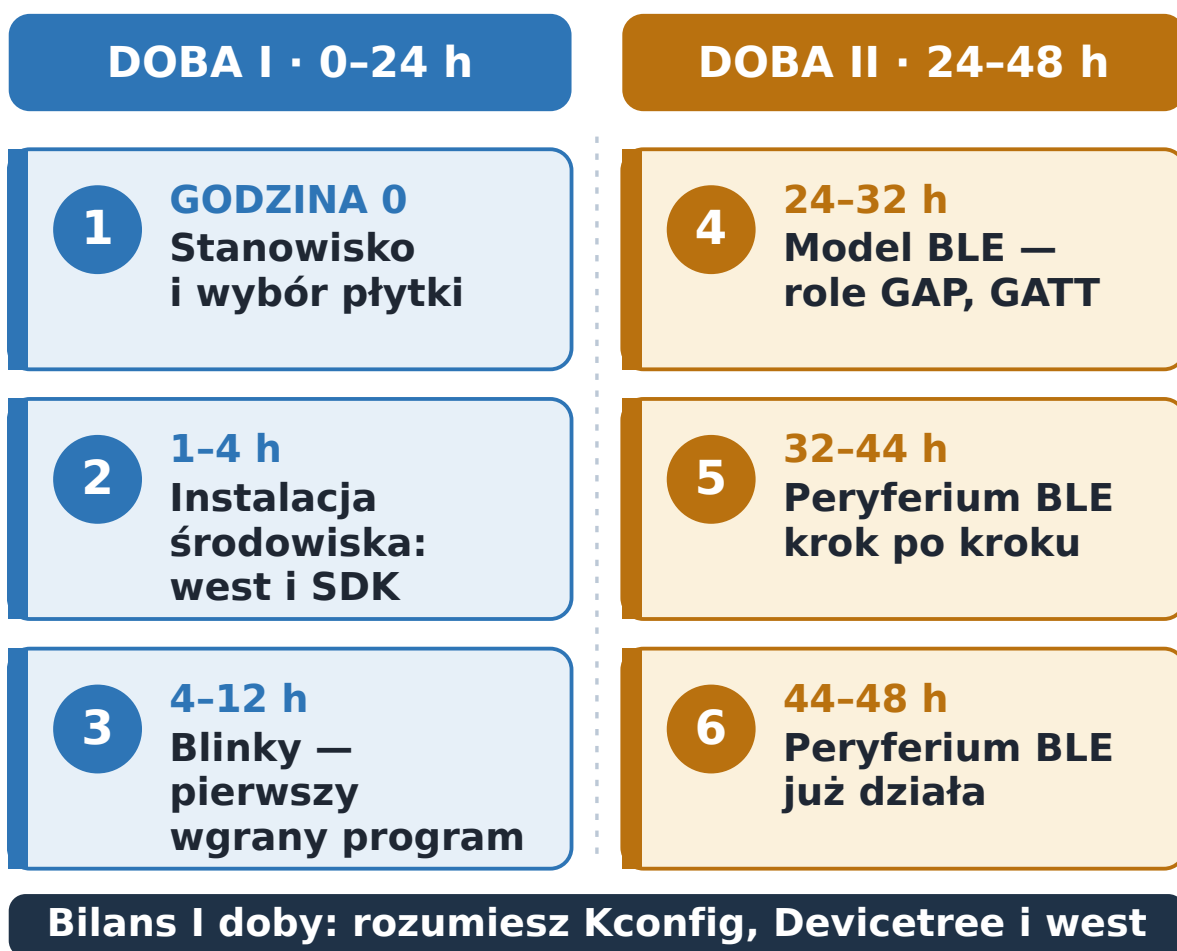
Ta wzmianka o wieloletniej eksploatacji nie jest marketingowym ozdobnikiem – tłumaczy, dlaczego po Zephyra sięgnęli najwięksi. Nordic Semiconductor, jeden z najwcześniejszych adoptujących i zarazem największy kontrybutor projektu, oparł na Zephyrze cały swój nRF Connect SDK. To istotna obserwacja praktyczna: ucząc się Zephyra na płytce Nordica, równocześnie uczymy się nRF Connect SDK. Wyboru Zephyra do swoich produktów dokonały także Google i Meta, a w gronie członkowskim projektu znajdziemy Intel, NXP czy producenta aparatów słuchowych Oticon. Dla magazynu takiego jak nasz to sygnał jednoznaczny: znajomość Zephyra przestała być ciekawostką, a stała się kompetencją zawodową.

Czym Zephyr różni się od RTOS-a, który większość z nas zna – od FreeRTOS? Trafną metaforę zaproponował Shawn Hymel, inżynier systemów wbudowanych i autor popularnej serii szkoleniowej „Introduction to Zephyr”, przygotowanej dla DigiKey. W jego ujęciu FreeRTOS jest biblioteką, a Zephyr – platformą. Z FreeRTOS-em to my decydujemy, jak poskładać wszystkie elementy w całość.

Z Zephyrem mamy gotową strukturę, która ma pomóc zapanować nad złożonością w skali większej niż pojedynczy projekt. Zephyr to bowiem nie tylko jądro i scheduler – to także kompletne stopy łączności (Bluetooth, Wi-Fi, sieć IP), setki pakietów wsparcia płytek, własny system budowania oraz metanarzędzia opisane w dalszej części artykułu.

Ta sama cecha, która czyni Zephyra potężnym, czyni go też wymagającym na starcie. Hymel nie owija tego w bawełnę: największym zarzutem wobec Zephyra jest krzywa uczenia się – i, jak pisze, krytycy nie są w błędzie. Zephyr potrafi być sporym wyzwaniem do podjęcia, zwłaszcza dla kogoś, kto przychodzi ze świata mikrokontrolerów i nie zetknął się wcześniej z Kconfigiem oraz Devicetree. W osobnym tekście Hymel mówi o tym jeszcze dosadniej: Zephyr ma realną złożoność, stromą krzywą uczenia się i nietrywialny koszt utrzymania, a koszty te nie zawsze są widoczne, gdy zaczynamy z systemem eksperymentować – tarcie ujawnia się później, gdy projekty i zespoły rosną.

Stąd pomysł na ten artykuł i na ramę „48 godzin” (rysunek 2). Pierwsza doba to instalacja środowiska i uruchomienie kanonicznego przykładu Blinky – migającej diody. Druga doba to droga od Blinky do pierwszego peryferium Bluetooth Low Energy. Taki podział nie jest sztuczny:



Rysunek 2. Mapa pierwszych 48 godzin z Zephyrem. Pierwsza doba kończy się nie migającą diodą, lecz zrozumieniem trzech filarów; druga prowadzi od modelu BLE do działającego peryferium

Tabela 1. Trzy popularne płytki widziane przez pryzmat 48-godzinnej drogi z Zephyrem. Ceny i dostępność celowo pominięto – kluczowa jest tu liczba etapów, które dana płytka obsługuje bez zmiany sprzętu

Płytki	Rdzeń i radio	Programowanie	Rola w drodze „od Blinky do BLE”
nRF52840-DK	Cortex-M4F, BLE 5.x na pokładzie	Wbudowany J-Link	Komplet – Blinky i BLE na jednej płytce; najlepiej udokumentowana
ESP32-S3-DevKitC	Xtensa LX7, BLE + Wi-Fi	Mostek USB-UART	BLE dostępne, lecz toolchain i flashowanie różni się od reszty artykułu
STM32 „Blue Pill”	Cortex-M3, bez radia	Zewnętrzny ST-Link/SWD	Świetna i tania do Blinky; do BLE wymaga osobnego modułu

odzwierciedla naturalną ścieżkę nauki, a co ważniejsze, koncentruje najgorsze tarcie dokładnie tutaj, na początku. Kto przejdzie te 48 godzin świadomie – z mapą, a nie po omacku – ten resztę Zephyra zwiedza już znacznie spokojniej.

2. Godzina zero: stanowisko pracy i wybór płytki referencyjnej

Zanim wpisujemy pierwszą komendę, dwie decyzje. Pierwsza dotyczy komputera-hosta. Zephyr oficjalnie wspiera Linuksa, macOS i Windowsa, a oficjalny przewodnik startowy prowadzi przez instalację na każdym z tych systemów. W praktyce najmniej problemów generuje Linux – natywnie lub w środowisku WSL2 na Windowsie. Toolchainy, narzędzia do flashowania i skrypty pomocnicze były projektowane przede wszystkim z myślą o nim. Jeśli więc mamy wybór, pierwsze 48 godzin warto przejść na Linuksie; jeśli nie – Windows również się sprawdzi, należy tylko liczyć się z odrobinę większą liczbą drobnych potknięć konfiguracyjnych.

Decyzja druga, ważniejsza dla tego artykułu, dotyczy płytki referencyjnej. Tytuł obiecuje drogę „od Blinky do BLE”, a to narzuca jeden warunek: ta sama płytka musi obsługiwać oba etapy. Dlatego jako płytkę referencyjną przyjmujemy w tym artykule nRF52840-DK firmy Nordic Semiconductor. Powodów jest kilka. Po pierwsze, układ nRF52840 ma na pokładzie radio Bluetooth Low Energy, więc droga do BLE nie wymaga zmiany sprzętu. Po drugie, płytka ma wbudowany programator SEGGER J-Link, co eliminuje osobny sprzęt do flashowania. Po trzecie – i najistotniejsze – to właśnie na nRF52840 oparta jest większość dostępnych kursów BLE w Zephyrze, w tym darmowa Nordic Developer Academy.

Wybierając tę płytkę, wchodzimy w najlepiej udokumentowany fragment ekosystemu.

Nie znaczy to, że inne płytki są złe – znaczy tylko, że dla tej konkretnej, 48-godzinnej drogi nRF52840-DK minimalizuje liczbę niespodzianek. Poniższe zestawienie pokazuje, jak na tym tle wypadają dwie inne popularne opcje.

Jest jeszcze jeden element wyposażenia stanowiska, o którym łatwo zapomnieć, a który okaże się niezbędny dopiero drugiej doby: skaner Bluetooth Low Energy. Bezpłatna aplikacja nRF Connect for Mobile (na Androida i iOS) pełni dla projektu BLE rolę, jaką dla projektu analogowego pełni oscyloskop – pozwala zobaczyć, czy urządzenie rozgłasza się poprawnie, połączyć się z nim i odczytać jego charakterystyki. Warto zainstalować ją już teraz, w godzinie zero, żeby nie szukać jej w pośpiechu w 30. godzinie. Przyda się też pakiet nRF Connect for Desktop z modułami Programmer i Serial Terminal.

3. Wyprawka kursanta: co przygotować przed startem

Skoro pierwsza doba zaczyna się od „godziny zero”, warto wiedzieć, co dokładnie powinno wtedy leżeć na biurku i być zainstalowane na dysku. Poniższe dwa zestawienia traktujemy jak listę zakupów i pobrań – kompletną wyprawkę, z którą da się przejść cały kurs bez przerw na doposażanie. Kurs prowadzimy ścieżką „czystego” Zephyra (metanarzędzie west i Zephyr SDK); osoby pracujące na płytkach Nordica mogą zamiast tego sięgnąć po nRF Connect SDK, który Zephyra w sobie zawiera.

Elementy materialne. Sprzęt, który trzeba mieć fizycznie pod ręką. Lista jest krótka – i celowo: do przejścia kursu wystarczy płytka, kabel i komputer.

Tabela 2. Wyposażenie materialne – sprzęt potrzebny do przejścia kursu. Wariant minimalny to płytka, kabel i komputer; resztę dobiera się zależnie od ambicji

Element	Rola w kursie	Skąd go wziąć
Płytki nRF52840-DK	Jedyny obowiązkowy zakup; obsługuje obie doby kursu – Blinky i BLE – bez zmiany sprzętu	Autoryzowani dystrybutorzy elektroniki (m.in. DigiKey, Mouser, Farnell, TME) Alternatywne płytki – patrz tabela 1
Kabel USB	Zasilanie płytki, wgrywanie programu i komunikacja z hostem (programator J-Link jest na płytce)	Zwykle dołączony do zestawu DK; w razie potrzeby dowolny kabel USB
Komputer (host)	Stanowisko pracy: kompilacja, wgrywanie i debugowanie	System Linux, macOS lub Windows; ok. 10...15 GB wolnego miejsca na dysku
Smartfon z Bluetooth LE	Rola centrali – skaner i klient w testach drugiej doby	Dowolny telefon z systemem Android lub iOS
Dioda LED, rezystor, przewody (opcjonalnie)	Tylko jeśli sygnał ma trafić poza płytke – podstawowa ścieżka korzysta z diod wbudowanych	Dowolny sklep z elektroniką; do ukończenia kursu zbędne

Tabela 3. Wyposażenie niematerialne – oprogramowanie, dane oraz wiedza i czas. Całość oprogramowania kursu jest otwarta i bezpłatna

Element	Rola w kursie	Skąd go wziąć
Zephyr SDK (toolchain)	Kompilatory i narzędzia dla architektur docelowych; od wydania 4.4 nosi nazwę Zephyr SDK 1.0	Wydania w repozytorium projektu na GitHubie (zephyrproject-rtos/sdk-ng); instalację opisuje rozdział 4
west oraz Python 3	Metanarzędzie zarządzające przestrzenią roboczą, budowaniem i wgrzywaniem	west instaluje się poleceniem pip; Python 3 – z python.org lub menedżera pakietów systemu
Git	Pobranie kodu źródłowego Zephyra i jego modułów	Strona git-scm.com lub menedżer pakietów systemu
Kod źródłowy Zephyra	Jądro RTOS, sterowniki, stos Bluetooth i katalog przykładów	Pobierany automatycznie poleceniami west init i west update (GitHub: zephyrproject-rtos/zephyr)
Edytor kodu lub IDE	Pisanie i przeglądanie kodu aplikacji	Dowolny; częsty wybór to bezpłatny Visual Studio Code
nRF Connect for Mobile	Skaner BLE na smartfonie – centrala w testach drugiej doby	Google Play lub App Store (wydawca: Nordic Semiconductor)
Symulatory BLE: Renode, BabbageSim (opcjonalnie)	Testowanie radia bez fizycznej płytki – przydatne, gdy sprzęt jeszcze nie dotarł (rozdział 9)	Projekty otwarte – renode.io oraz repozytoria środowiska bsim
Wiedza wstępna	Język C, podstawy pracy w wierszu poleceń, ogólne pojęcie o mikrokontrolerach	Kurs zakłada tę wiedzę i jej nie naucza – warto uzupełnić ją przed startem
Czas	Rama kursu: dwie doby pracy z marginesem na eksperymenty	Ok. 48 godzin, do rozłożenia we własnym tempie

KOSZT, KONTA, CZAS

Jedynym realnym wydatkiem jest płytka rozwojowa – całe oprogramowanie kursu (Zephyr SDK, west, kod źródłowy Zephyra, nRF Connect for Mobile) jest otwarte i darmowe. Do przejścia kursu nie trzeba zakładać żadnego konta; opcjonalne jest jedynie konto w Nordic Developer Academy, jeśli zechcemy korzystać z jej materiałów. Budżet czasu to tytułowe 48 godzin – rozłożone tak, jak pozwala kalendarz.

Elementy niematerialne. Oprogramowanie, dane oraz – równie ważne – wiedza i czas. Tę część w całości da się skompletować bezpłatnie; rozdział 4 przeprowadza przez samą instalację krok po kroku.

4. Pierwsze godziny: instalacja środowiska

Instalacja Zephyra dezorientuje nowicjuszy, bo nie jest jedną instalacją, lecz trzema. Warto mieć ten model z tyłu głowy od początku, bo porządkuje cały proces.

- **west i zależności Pythona** – metanarzędzie projektu, opisane szerzej w rozdziale 6, instalowane przez menedżer pakietów Pythona pip, najlepiej wewnątrz wirtualnego środowiska.
- **przestrzeń robocza (workspace)** – katalog roboczy zawierający drzewo źródeł Zephyra wraz

z dziesiątkami modułów. To największy objętościowo element – liczony w gigabajtach – i jego pobranie zajmuje najwięcej czasu.

- **Zephyr SDK** – osobny pakiet zawierający kompilatory, asemblyery i linkery dla wszystkich wspieranych architektur. Od wydania Zephyr 4.4 (kwiecień 2026) obowiązuje Zephyr SDK w wersji 1.0.

Sekwencja komend wygląda w uproszczeniu tak, jak na **listingu 1**. Każdy wiersz to osobny krok oficjalnego przewodnika startowego; szczegóły zależne od systemu operacyjnego – w tym instalacja samego SDK – opisane są w dokumentacji i w praktycznym przewodniku Bootlin.

Dwie uwagi praktyczne na tym etapie. Po pierwsze, wirtualne środowisko Pythona trzeba aktywować za każdym razem, gdy wracamy do pracy – to częste źródło zagadkowych błędów u początkujących, którzy o tym zapominają.

```
# 1. Wirtualne środowisko Pythona i metanarzędzie west
python -m venv ~/zephyrproject/.venv
source ~/zephyrproject/.venv/bin/activate
pip install west

# 2. Pobranie źródeł Zephyra i modułów (najdłuższy krok)
west init ~/zephyrproject
cd ~/zephyrproject
west update

# 3. Eksport pakietu CMake i doinstalowanie zależności
west zephyr-export
west packages pip --install

# 4. Instalacja Zephyr SDK (toolchainy) – wg dokumentacji
```

Listing 1. Szkielet instalacji środowiska Zephyra. Polecenie west update pobiera kilka gigabajtów danych – to normalne; dobrym momentem na kawę jest właśnie ten krok

```
cd ~/zephyrproject/zephyr
west build -p always -b nrf52840dk/nrf52840 samples/basic/blink
west flash
```

Listing 2. Budowanie i flashowanie przykładu Blinky. Flaga `-p always` wymusza czystą kompilację, a identyfikator `nrf52840dk/nrf52840` to nazwa płytki w nowym modelu opisu sprzętu Zephyra

```
/* Dioda opisana w Devicetree aliasem "led0" */
static const struct gpio_dt_spec led =
    GPIO_DT_SPEC_GET(DT_ALIAS(led0), gpios);

int main(void)
{
    gpio_pin_configure_dt(&led, GPIO_OUTPUT_ACTIVE);

    while (1) {
        gpio_pin_toggle_dt(&led);
        k_msleep(500);
    }
    return 0;
}
```

Listing 3. Trzon przykładu Blinky w skróconej, poglądowej formie. Kod nie wymienia żadnego konkretnego pinu ani portu GPIO

Po drugie, krok `west update` bywa frustrujący nie dlatego, że coś jest zepsute, lecz dlatego, że trwa długo i pobiera dużo. To nie awaria – to skala projektu. Kto wie o tym wcześniej, ten nie traci czasu na diagnozowanie nieistniejącego problemu. Po przejściu tych czterech kroków mamy kompletne, działające środowisko – i możemy zbudować pierwszy program.

5. Blinky: anatomia najprostszego projektu Zephyra

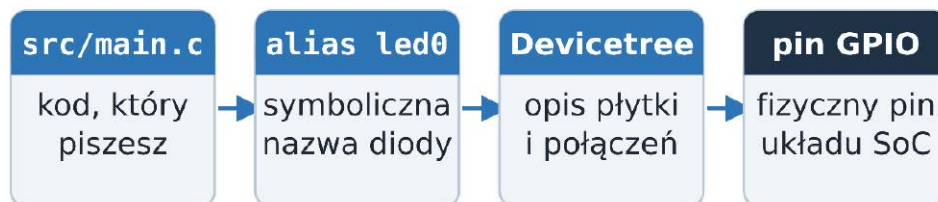
Blinky – program migający diodą – jest w świecie systemów wbudowanych odpowiednikiem „Hello, world”.

W Zephyrze znajdziemy go w repozytorium projektu, w katalogu `samples/basic/blink`. Zbudowanie i wgranie go na płytkę referencyjną sprowadza się do dwóch poleceń (listing 2).

Jeśli dioda LED1 zaczęła migać – gratulacje, pierwsza połowa drogi za nami w sensie wykonawczym. Teraz najważniejsze: zrozumieć, co właściwie się wydarzyło. Aplikacja Zephyra to nie pojedynczy plik `.c`. To katalog, którego trzon stanowią cztery elementy: plik `CMakeLists.txt` (spina projekt z systemem budowania), plik `prj.conf` (konfiguracja oprogramowania – o tym w rozdziale 6), opcjonalna nakładka Devicetree



Jak Blinky znajduje diodę bez podawania pinu



Ten sam plik `main.c` buduje się na setkach płytek Zephyra

Rysunek 3. Anatomia aplikacji Zephyra. U góry – pliki tworzące projekt; u dołu – łańcuch, dzięki któremu przykład Blinky trafia do właściwego pinu bez wpisywania jego numeru w kodzie

Tabela 4. Ten sam efekt – migająca dioda – w dwóch światach. Różnica nie polega na liczbie linii kodu, lecz na tym, gdzie zapisana jest wiedza o sprzęcie

Aspekt	Blinky bare-metal	Blinky w Zephyrze
Wskazanie diody	Adres rejestru konkretnego pinu	Alias led0 w Devicetree
Zmiana płytki	Przepisanie kodu dostępu do GPIO	Zmiana parametru -b w poleceniu budowania
Opóźnienie 500 ms	Pętla zwłoki lub własny timer	k_msleep() – usługa jądra RTOS
Konfiguracja funkcji	Dyrektywy #define i ifdef w kodzie	Plik prj.conf (Kconfig)

Tabela 5. Trzy filary, które trzeba zrozumieć w pierwszej dobie. Dwa z nich Zephyr przejął wprost z ekosystemu Linuksa – to nie przypadek, lecz świadoma decyzja projektowa

Mechanizm	Odpowiada na pytanie	Gdzie żyje	Skąd pochodzi
Kconfig	Jakie funkcje oprogramowania włączyć?	prj.conf, symbole CONFIG_*	Z jądra Linuksa
Devicetree	Jaki sprzęt istnieje i jak jest podłączony?	Pliki .dts oraz .overlay	Z jądra Linuksa
west	Jak zarządzać workspace i całym cyklem pracy?	Polecenia west w terminalu	Narzędzie własne Zephyra

z rozszerzeniem .overlay oraz katalog src/ z kodem źródłowym.

Sednem Blinky jest jednak coś, co na pierwszy rzut oka umyka. Przyjrzyjmy się fragmentowi kodu na **listingu 3**.

To jest właśnie clou. W programie bare-metal na NRF52840 zapalilibyśmy diodę, odwołując się wprost do rejestru konkretnego pinu konkretnego portu. Kod Blinky nie robi tego nigdzie. Zamiast tego pyta o alias led0 – symboliczną nazwę zdefiniowaną w opisie sprzętu (Devicetree). To, który fizyczny pin kryje się pod tym aliasem, wie płytka, nie program. Dzięki temu dokładnie ten sam, niezmienny plik źródłowy Blinky kompiluje się i działa na setkach różnych płytek wspieranych przez Zephyra. Przenośność nie jest tu dodatkiem – jest fundamentem architektury (rysunek 3).

6. Trzy pojęcia, o które potkniesz się najszybciej: Kconfig, Devicetree i west

Jeśli pierwsza doba z Zephyrem ma jeden prawdziwy próg, to jest nim właśnie ten rozdział. Shawn Hymel wskazuje go imiennie: trudność Zephyra dla osoby ze świata mikrokontrolerów bierze się przede wszystkim z niezajomości Kconfiga i Devicetree. Dobra wiadomość jest taka, że to tylko trzy pojęcia – i że każde z nich ma jedno, dające się wyrazić w jednym zdaniu zadanie.

Kconfig odpowiada za oprogramowanie. To system konfiguracji decydujący o tym, które elementy Zephyra zostaną w ogóle skompilowane do naszego obrazu – czy ma być stos Bluetooth, czy logowanie, czy obsługa konkretnego sterownika. Włącza się je symbolami postaci CONFIG_* w pliku prj.conf. To dokładnie ten sam Kconfig, który od lat konfiguruje jądro Linuksa. Można go też przeglądać interaktywnie poleceniem west build -t menuconfig.

Devicetree opisuje sprzęt. To formalny opis tego, jaki sprzęt znajduje się w systemie i jak jest połączony – procesor, jego peryferia, diody, magistrale, czujniki. Płytkę ma swój plik .dts opisujący układ i połączenia; my, jeśli trzeba, dopisujemy własną nakładkę .overlay, która ten opis

modyfikuje. To stąd Blinky z poprzedniego rozdziału wiedział, czym jest led0.

west spina całość. To metanarzędzie zarządzające przestrzenią roboczą i obejmujące cały cykl pracy – od pobrania źródeł (west update), przez budowanie (west build), po wgrywanie i debugowanie (west flash, west debug). Jeśli Kconfig i Devicetree to dwie mapy, west jest narzędziem, którym się po nich poruszamy.

Kluczowy moment pierwszej doby to chwila, w której te trzy elementy układają się w jedną myśl: Zephyr konsekwentnie rozdziela trzy pytania, które w programowaniu bare-metal zlewają się w jedno. Pytanie „jaki mam sprzęt” należy do Devicetree. Pytanie „jakiego oprogramowania chcę” należy do Kconfiga. Pytanie „jak to wszystko zbudować i wgrać” należy do west (**rysunek 4**). Gdy ten podział raz zaskoczy, reszta Zephyra przestaje być labiryntem i staje się czymś, po czym da się nawigować. To jest właściwy zysk pierwszych dwudziestu czterech godzin – nie migająca dioda, lecz ten sposób myślenia.

7. Druga doba: jak Zephyr widzi Bluetooth Low Energy

Bluetooth Low Energy ma własną, zasłużoną reputację technologii o stromej krzywej uczenia się – niezależnie od RTOS-a. Mohammad Afaneh, założyciel serwisu Novel Bits i Bluetooth Developer Academy, autor książki „Intro to Bluetooth Low Energy” i konsultant współpracujący między innymi z organizacją Bluetooth SIG, otwarcie opisuje własne początki: gdy zaczynał poznawać BLE od strony oprogramowania wbudowanego, spędził godziny, dni i tygodnie na lekturze wszystkiego, co znalazł, krzywa uczenia się była stroma, a przedzieranie się przez specyfikację Bluetooth – jak to ujmuję – wręcz frustrujące. Dobra wiadomość brzmi tak: znaczną część tej specyfikacji Zephyr bierze na siebie.

Zephyr dostarcza bowiem kompletny stos Bluetooth Low Energy – zarówno warstwę kontrolera (Controller), jak i hosta (Host) oraz łączące je HCI (**rysunek 5**).



Rysunek 4. Trzy filary, na których stoi każda aplikacja Zephyra. Każdy odpowiada na inne pytanie; dwa z nich – Kconfig i Devicetree – Zephyr przejął wprost z ekosystemu jądra Linuksa

BILANS PIERWSZEJ DOBY

Po 24 godzinach mamy: działające środowisko (west, workspace, SDK), zbudowany i wgrany przykład Blinky oraz – co najważniejsze – model myślowy trzech filarów. Potrafimy odpowiedzieć, gdzie opisuje się sprzęt, gdzie włącza funkcje oprogramowania i czym buduje się projekt. To wystarczający fundament, by drugą dobę poświęcić już nie na walkę z narzędziami, lecz na rzecz właściwą – radio.

APLIKACJA

Twój serwis GATT i logika

HOST

GAP · GATT · ATT · L2CAP · SMP

HCI

Host Controller Interface

KONTROLER

Link Layer, obsługa czasu radiowego

RADIO 2,4 GHz

sprzęt radiowy układu SoC

- Warstwa aplikacji — kod, który piszesz (serwis GATT)
- Stos Bluetooth Low Energy dostarczany przez Zephyr
- Warstwa sprzętowa — radio 2,4 GHz w układzie SoC

Rysunek 5. Uproszczony model stosu Bluetooth Low Energy w Zephyrze. Warstwy kontrolera, HCI i hosta dostarcza system – projektant pisze wyłącznie warstwę aplikacji

To nie jest stos zabawkowy. Jego jakość bierze się stąd, że utrzymują go zawodowi specjaliści: podsystemem Bluetooth w Zephyrze opiekuje się Carles Cufí z Nordic Semiconductor – inżynier z ponad ćwierćwieczem doświadczenia w oprogramowaniu układowym, współautor wydanej przez O'Reilly książki „Getting Started with Bluetooth Low Energy” i wieloletni menedżer wydań Zephyra. Skala zaufania do tego stosu jest mierzalna: Nordic, którego portfolio na nim się opiera, deklaruje wysyłkę ponad miliona układów BLE dziennie.

Zanim napiszemy pierwszą linię kodu radiowego, warto opanować pięć pojęć, które tworzą model BLE – i które Zephyr odwzorowuje niemal jeden do jednego. Role GAP określają, czym urządzenie jest w eterze: peryferium (peripheral) i centralą (central) dla połączeń, rozgłaszającym (broadcaster) i obserwatorem (observer) dla komunikacji bezpołączeniowej. Rozgłaszanie (advertising) to nadawanie krótkich pakietów, dzięki którym urządzenie jest w ogóle widoczne. Połączenie (connection) to nawiązana, dwukierunkowa relacja między peryferium a centralą.



Rysunek 6. Budowa peryferium BLE w trzech krokach. Trudność i zakres kontroli nad protokołem rosną z każdym krokiem – od prostego rozgłaszania po własny serwis GATT

GATT to model danych: serwisy grupujące powiązane informacje oraz charakterystyki – pojedyncze, odczytywalne i zapisywalne wartości. Profil to wreszcie uzgodniony zestaw serwisów realizujący konkretną funkcję.

Tytułowe przejście „od GPIO do radia” ma tu głębszy sens, niż się wydaje. Pierwszej doby naszym peryferium był pojedynczy pin z diodą. Drugiej doby „peryferium” to całe nasze urządzenie, widziane z eteru przez telefon czy inny sprzęt. Zmienia się skala, nie sposób myślenia – i, co istotne, nie zmieniają się narzędzia. Włączenie Bluetootha to znów Kconfig: wystarczą symbole takie jak `CONFIG_BT` i `CONFIG_BT_PERIPHERAL` w pliku `prj.conf`. Cała wiedza o trzech filarach z pierwszej doby przenosi się na drugą bez żadnej straty.

8. Peryferium BLE krok po kroku – od reklamy do własnego serwisu GATT

Drogę do pierwszego peryferium najrozsądniej pokonać w trzech krokach o rosnącej trudności (rysunek 6), za każdym razem opierając się na oficjalnym przykładzie z katalogu `samples/bluetooth` – nie pisząc kodu

od zera. To powtarzalna rada wszystkich cytowanych tu praktyków: gotowe przykłady Zephyra są wzorcem do naśladowania, nie tylko materiałem do oglądania.

Krok pierwszy – samo rozgłaszanie. Najprostsze sensowne urządzenie BLE nie nawiązuje połączeń; jedynie się rozgłasza, jak beacon. Wystarczy zainicjować stos wywołaniem `bt_enable()` i uruchomić rozgłaszanie z przygotowaną tablicą danych. Od tej chwili urządzenie jest widoczne w aplikacji nRF Connect for Mobile – i to jest pierwszy namacalny sukces drugiej doby.

Krok drugi – peryferium z gotowym serwisem. Następnym etapem jest urządzenie, z którym da się połączyć i które udostępnia standardowy serwis. Najwładźniejszy jest serwis baterii (Battery Service, BAS). Tu Zephyr pokazuje swoją siłę: cały serwis dostajemy gotowy, włączając go jednym symbolem Kconfig – `CONFIG_BT_BAS`, widocznym już na listingu 4. Nie implementujemy protokołu; jedynie aktualizujemy poziom naładowania, a obsługą GATT zajmuje się stos.

Krok trzeci – własny serwis GATT. Dopiero teraz, mając pewny grunt, definiujemy serwis autorski – z własnym 128-bitowym identyfikatorem UUID i własną

```
# prj.conf – włączenie roli peryferium i serwisu baterii
CONFIG_BT=y
CONFIG_BT_PERIPHERAL=y
CONFIG_BT_DEVICE_NAME="EP-Peripheral"
CONFIG_BT_BAS=y
```

Listing 4. Konfiguracja Kconfig dla peryferium BLE z gotowym serwisem baterii. Cztery linie – i to one, a nie kod, decydują o tym, że w obrazie znajdzie się stos Bluetooth


```
static const struct bt_data ad[] = {
    BT_DATA_BYTES(BT_DATA_FLAGS, BT_LE_AD_GENERAL),
    BT_DATA(BT_DATA_NAME_COMPLETE, "EP-Peripheral", 13),
};

int main(void)
{
    bt_enable(NULL);
    bt_le_adv_start(/* parametry */, ad, ARRAY_SIZE(ad),
        NULL, 0);
    /* urządzenie jest teraz widoczne w skanerze BLE */
    return 0;
}
```

Listing 5. Poglądowy szkielec rozgłaszania. Konkretne nazwy makr – zwłaszcza parametrów rozgłaszania – zmieniają się między wydaniem Zephyra; aktualne wartości należy każdorazowo potwierdzić w dokumentacji danej wersji

charakterystyką, na przykład sterującą diodą po radiu (klasyczne ćwiczenie „Smart LED” opisywane w przewodnikach Michaela Angerera i zespołu Hubble). Służy do tego makro `BT_GATT_SERVICE_DEFINE`, w którym deklarujemy charakterystyki oraz ich uprawnienia – odczyt, zapis, powiadomienia. To jest moment, w którym przestajemy korzystać z cudzych serwisów, a zaczynamy projektować własny protokół aplikacyjny.

Na każdym z trzech kroków narzędziem weryfikacji jest wspomniana aplikacja `nRF Connect for Mobile`, pełniąca rolę centrali: skanuje eter, łączy się z naszym peryferium, odczytuje i zapisuje charakterystyki. Bez niej pracowalibyśmy na ślepo – dlatego właśnie warto było zainstalować ją już w godzinie zero.

9. Ostatnie godziny: BLE bez płytki i mapa dalszej drogi

Co, jeśli płytka jeszcze nie dotarła albo chcemy eksperymentować w trasie? Zephyr daje na to odpowiedź, której klasyczny RTOS nie ma. Stos Bluetooth można uruchamiać i testować w symulacji – bez fizycznego radia. Projekt utrzymuje w tym celu środowisko `BabbleSim` (`bsim`), służące do symulowania komunikacji BLE i wykorzystywane w testach regresyjnych samego Zephyra. Równoległe środowisko `Renode` potrafi uruchamiać przykłady BLE Zephyra w symulowanym sprzęcie. Dla osoby uczącej się oznacza to, że ostatnie godziny naszych 48 można spędzić na eksperymentach z radiem, nawet nie mając go pod ręką.

Czterdziesta ósma godzina to jednak nie meta, lecz punkt orientacyjny. Mamy za sobą działające Blinky i działające peryferium BLE – i, co ważniejsze, rozumiemy, dlaczego działają. Naturalne kolejne kroki to: zarządzanie energią (kluczowe dla realnych urządzeń BLE), bezpieczna aktualizacja firmware'u przez bootloader `MCUboot`, ujednolicony interfejs czujników Zephyra oraz łączność z chmurą. Każdy z tych tematów to materiał na osobny artykuł.

Dokąd po wiedzy? Wśród źródeł sprawdzonych przez praktyków warto wymienić oficjalną dokumentację

Zephyra i jego przewodnik startowy; darmową, ustrukturyzowaną `Nordic Developer Academy` z kursem podstaw BLE; serię „Introduction to Zephyr” Shawna Hymela przygotowaną dla DigiKey; materiały firmy `Golioth` oraz utrzymywaną przez nią kuratorską listę `awesome-zephyr-rtos`, zbierającą artykuły, przykłady i narzędzia. To nie przypadek, że autorzy tych zasobów to w dużej mierze ci sami eksperci, których opinie towarzyszą nam w tym artykule.

10. Wskazówki praktyków: jak przejść te 48 godzin bez frustracji

Na koniec zbierzmy w jedno to, co o pierwszym kontakcie z Zephyrem radzą cytowani tu specjaliści. Te wskazówki nie skrócą drogi – ale sprawią, że jej najtrudniejszy odcinek przejdziemy świadomie.

- **Pogódź się z tym, że pierwsze starcie boli.** Shawn Hymel mówi o tym wprost: krzywa uczenia się Zephyra jest stroma, a tarcie na starcie jest realne. To nie znak, że robimy coś źle – to wpisany w platformę koszt wejścia. Kto się tego spodziewa, ten nie zniechęca się w trzeciej godzinie.
- **Ucz się jednego filaru naraz.** Nie próbuj opanować `Kconfig`, `Devicetree` i `west` jednocześnie. Pierwsza doba to po to, by każdemu z trzech filarów dać osobną uwagę. Gdy zrozumiesz, które pytanie należy do którego mechanizmu, reszta układu się sama.
- **Startuj z gotowych przykładów.** Oficjalne przykłady z katalogów `samples/basic` i `samples/bluetooth` to nie demonstracje do oglądania, lecz wzorce do naśladowania. Każde sensowne peryferium BLE zaczynać od skopiowania działającego przykładu, a nie od pustego pliku.
- **Trzymaj się jednej płytki.** Tytuł tego artykułu – „od Blinky do BLE” – działa tylko wtedy, gdy oba etapy realizujemy na jednym sprzęcie. Wybierz jedną płytkę referencyjną i nie zmieniaj jej przez całe 48 godzin.
- **Przygotuj narzędzia obserwacji wcześniej.** Skaner BLE w telefonie jest dla projektu radiowego tym, czym

oscilloskop dla analogowego. Miej go zainstalowany, zanim napiszesz pierwszą linię kodu Bluetootha – inaczej będziesz pracować po omacku.

- **Pamiętaj, po co to robisz.** Hymel formułuje to jako strategiczną radę: dla pojedynczego, małego prototypu narzut Zephyra potrafi wydać się ciężki, ale dla rodziny produktów czy systemu, który ma się skalować, ta sama struktura staje się realnym atutem. Wiedząc, do czego Zephyr jest dobry, łatwiej znieść jego koszt wejścia.

I myśl spinająca, znów od Kate Stewart: następna dekada Zephyra – jak mówi – będzie definiowana nie tyle przez nowe funkcje, ile przez wierność zasadom, które uczyniły projekt skutecznym: otwartości, współpracy, przenośności i zaufaniu. Pierwsze 48 godzin to mała inwestycja w wejście do ekosystemu zbudowanego dokładnie na tych zasadach. Migająca dioda i rozgłaszające się peryferium BLE to dopiero przedsmak – ale to przedsmak, który otwiera drzwi do reszty.

JS, JT

Pierwsze 48 godzin z Zephyrem nie czyni z nikogo eksperta – i nie taki jest ich cel. Ich celem jest przekroczenie progu: przejście od stanu, w którym Zephyr jest niezrozumiałym labiryntem narzędzi, do stanu, w którym jest mapą, po której potrafimy się poruszać. Kto ten próg przekroczy świadomie, dla tego każda kolejna godzina z Zephyrem jest już znacznie łatwiejsza niż pierwsza.

Źródła

- [1] Zephyr Project – „Getting Started Guide”. https://docs.zephyrproject.org/latest/develop/getting_started/index.html
- [2] Zephyr Project – strona projektu i materiały wprowadzające. <https://www.zephyrproject.org>
- [3] „Zephyr Turns 10 as Global Adoption Surges and Long-Term Embedded Use Expands” – komunikat Zephyr Project/Linux Foundation, marzec 2026.
<https://www.zephyrproject.org/zephyr-turns-10-as-global-adoption-surges-and-long-term-embedded-use-expands/>
- [4] „Zephyr RTOS 4.4 Now Available: WireGuard, Wi-Fi Direct, OpenRISC and More” – blog Zephyr Project, kwiecień 2026.
<https://www.zephyrproject.org/zephyr-rtos-4-4-now-available-wireguard-wi-fi-direct-openrisc-and-more/>
- [5] Bootlin – „Getting started with Zephyr”. <https://bootlin.com/blog/getting-started-with-zephyr/>
- [6] S. Hymel, „Why Use Zephyr? A Practical Guide for Embedded Engineers Choosing the Right RTOS”, shawnhymel.com, 2025.
<https://shawnhymel.com/2741/why-use-zephyr-a-practical-guide-for-embedded-engineers-choosing-the-right-rtos/>
- [7] S. Hymel, „Zephyr vs FreeRTOS: How to Choose the Right RTOS for Your Embedded Project”, shawnhymel.com, 2025.
<https://shawnhymel.com/3106/zephyr-vs-freertos-how-to-choose-the-right-rtos-for-your-embedded-project/>
- [8] S. Hymel, „The Hidden Costs of Using Zephyr (and How to Mitigate Them)”, shawnhymel.com, luty 2026.
<https://shawnhymel.com/3193/the-hidden-costs-of-using-zephyr-and-how-to-mitigate-them/>
- [9] S. Hymel, „Introduction to Zephyr” – 12-częściowa seria szkoleniowa DigiKey/Maker.io (część 1).
<https://www.digikey.com/en/maker/tutorials/2025/introduction-to-zephyr-part-1-getting-started-installation-and-blink>
- [10] K. Townsend, C. Cufi, Akiba, R. Davidson, „Getting Started with Bluetooth Low Energy”, O'Reilly Media, 2014.
<https://www.oreilly.com/library/view/getting-started-with/9781491900550/>
- [11] M. Afaneh, „Intro to Bluetooth Low Energy”, Novel Bits. <https://novelbits.io/intro-to-ble-book/>
- [12] Novel Bits – „Zephyr Tutorial: Bluetooth Low Energy Development”.
<https://novelbits.io/zephyr-getting-started-bluetooth-low-energy-development/>
- [13] Hubble – „Intro to Zephyr RTOS: Build a BLE Peripheral from Scratch”.
<https://hubble.com/community/guides/intro-to-zephyr-rtos-build-a-ble-peripheral-from-scratch-with-code-examples/>
- [14] M. Angerer – „Zephyr Basics: Bluetooth Low Energy (BLE) – Part 1: Peripheral”.
<https://michaelangerer.dev/zephyr/2022/04/07/zephyr-basics-ble-1.html>
- [15] Zephyr Project – „Developing Bluetooth Low Energy applications with Zephyr RTOS: Dissecting an iBeacon example” (K. Vervloesem).
<https://www.zephyrproject.org/developing-bluetooth-low-energy-apps-with-ibeacon-ex/>
- [16] Renode – „Bluetooth Low Energy simulation” (dokumentacja); demonstracja rozwijana na Zephyrze i układzie nRF52840.
<https://renode.readthedocs.io/en/latest/tutorials/ble-simulation.html>
- [17] Holisticon – „A practical guide to Zephyr OS”. <https://holisticon.pl/blog/zephyr-os-guide-embedded-devices-real-time-solutions/>
- [18] goliath/awesome-zephyr-rtos – kuratorska lista zasobów (GitHub). <https://github.com/goliath/awesome-zephyr-rtos>
- [19] Nordic Semiconductor – Nordic Developer Academy, kurs „Bluetooth Low Energy Fundamentals”.
<https://academy.nordicsemi.com/courses/bluetooth-low-energy-fundamentals/>
- [20] zephyrproject-rtos/zephyr – repozytorium projektu; przykład samples/basic/blinky (GitHub).
<https://github.com/zephyrproject-rtos/zephyr/tree/main/samples/basic/blinky>

REKLAMA

Publikujemy dla projektantów i programistów elektroniki

ELPORTAL.pl

Półprzewodniki szerokoprzerwowe

GaN vs. SiC – który wybrać do przetwornicy 200 W?

Granica decyzyjna przebiega przy 650 V – i właśnie tam, w klasie mocy rzędu 200 W, azotek galu i węglík krzemu spotykają się w bezpośrednim starciu. Zanim porównamy obu rywali, przypominamy, co w ogóle łączy te dwa materiały i dlaczego oba wypierają dziś krzem. Przedstawimy, co o wyborze technologii mówią eksperci oraz konstruktorzy, którzy obie technologie porównali w praktyce.

„GaN czy SiC?” – to pytanie pojawia się dziś przy projektowaniu niemal każdej przetwornicy średniej mocy. Przez lata odpowiedź była prosta, bo obie technologie zajmowały rozłączne nisze. Węglík krzemu (SiC) królował tam, gdzie liczyły się wysokie napięcia i duże moce; azotek galu (GaN) – tam, gdzie potrzebne były ekstremalne częstotliwości i miniaturyzacja. Domeny się nie nakładały, więc wybór nie przedstawiał problemu.

Ta sytuacja się skończyła. Na rynku dostępne są dziś zarówno tranzystory GaN HEMT, jak i MOSFET-y SiC o napięciu znamionowym 650 V – i właśnie w tym punkcie obie technologie zderzyły się czołowo. Przetwornica 200 W zasilana z szyny rzędu 400 V mieści się dokładnie w strefie nakładania: napięcie jest na tyle niskie, że GaN ma pełne prawo startu, a moc – na tyle umiarkowana, że SiC nie ma jeszcze przewagi, którą zyskuje przy kilku kilowatach. Dla konstruktora pytanie „GaN czy SiC?” przestało być formalnością, a stało się realną, ważną decyzją projektową.

W tym artykule zbieramy to, co o owym wyborze mówią eksperci – uczestnicy branżowych debat, autorzy not aplikacyjnych czołowych producentów oraz konstruktorzy, którzy obie technologie postawili obok siebie na stanowisku pomiarowym. Zaczniemy jednak od kroku wstecz: od tego, co GaN i SiC w ogóle łączy i dlaczego oba materiały coraz śmielej zastępują krzem.

Półprzewodniki szerokoprzerwowe: co krzem oddaje nowym materiałom

Przez kilkadziesiąt lat krzem był materiałem bezalternatywnym. Tranzystory MOSFET i IGBT na krzemie zbudowały całą współczesną elektronikę mocy – od zasilaczy komputerowych po napędy przemysłowe. Problem w tym,

że krzem jako materiał ma swoje granice fizyczne, a w wielu zastosowaniach konstruktorzy zaczęli się o nie wprost opierać. Aby przyrząd krzemowy zablokował wysokie napięcie, musi mieć odpowiednio grubą i słabo domieszkowaną warstwę dryfu – a to oznacza wysoką rezystancję w stanie włączenia i duże straty przewodzenia. Dalsza poprawa parametrów przestała być kwestią lepszego procesu; stała się kwestią materiału.

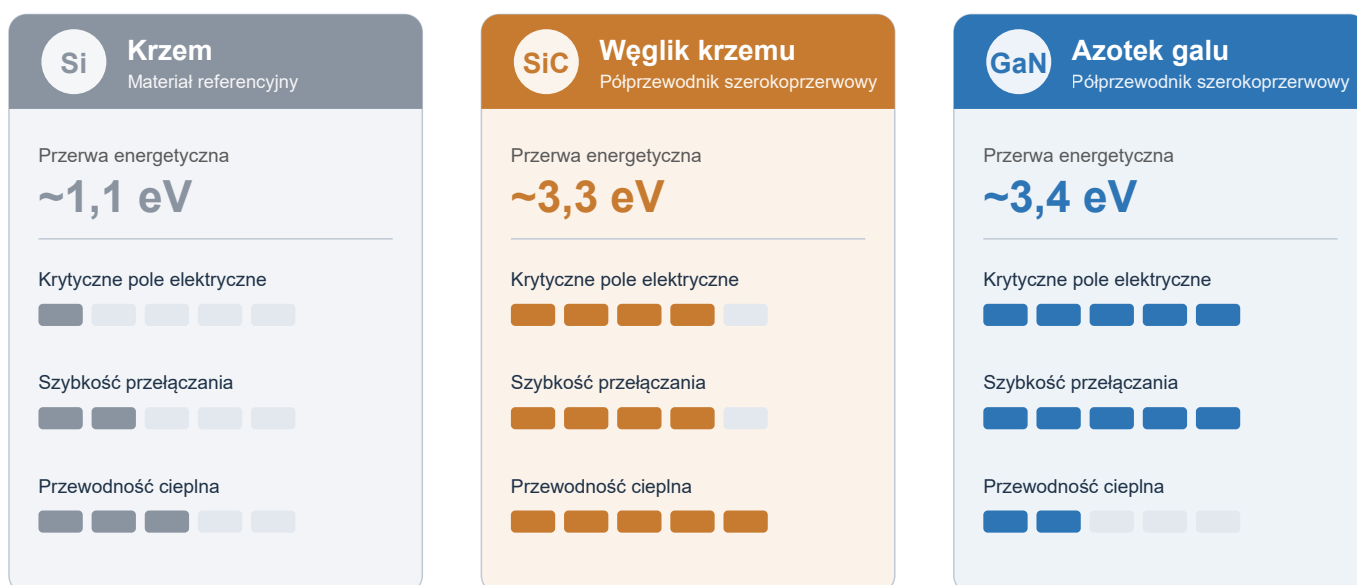
Tu wkraczają półprzewodniki szerokoprzerwowe, oznaczane skrótem WBG (od ang. wide-bandgap). Nazwa odnosi się do przerwy energetycznej – odstęp między pasmem walencyjnym a pasmem przewodnictwa. W krzemie wynosi ona około 1,1 eV. W węglíku krzemu (4H-SiC) to już mniej więcej 3,3 eV, a w azotku galu – około 3,4 eV, czyli blisko trzykrotnie więcej niż w krzemie. Sama liczba mówi niewiele, ale jej konsekwencje są dla elektroniki mocy fundamentalne.

Najważniejszą z nich jest krytyczne pole elektryczne – natężenie pola, przy którym materiał ulega przebiciu. Dla SiC i GaN jest ono blisko dziesięciokrotnie wyższe niż dla krzemu. W praktyce oznacza to, że tę samą wartość napięcia blokującego można uzyskać w warstwie dryfu około dziesięć razy cieńszej. Cieńsza warstwa to znacznie niższa rezystancja w stanie włączenia przy tej samej powierzchni struktury – a więc dramatycznie mniejsze straty przewodzenia, albo ten sam przyrząd o znacznie wyższym napięciu znamionowym w tej samej obudowie. To właśnie ta właściwość materiału otwiera elektronice mocy nowe możliwości.

Druga konsekwencja dotyczy szybkości. Nośniki ładunku w przyrządach WBG – w szczególności w azotku galu, gdzie prąd płynie w tak zwanym dwuwymiarowym gazie elektronowym – przemieszczają się szybko, a pojemności

Co zmienia szeroka przerwa energetyczna

Krzem a półprzewodniki szerokoprzerwowe — poglądowe porównanie kluczowych właściwości materiału



Wskaźniki mają charakter poglądowy i ilustrują relacje między materiałami, a nie wartości katalogowe konkretnych przyrządów.

Rysunek 1. Krzem a półprzewodniki szerokoprzerwowe – poglądowe zestawienie kluczowych właściwości materiałów

pasówytnicze, które trzeba przy każdym przełączeniu naładować i rozładować, są niewielkie. Przyrząd WBG przełączasięwciążznacznie szybciej niż jego krzemowy odpowiednik i może pracować przy wyraźnie wyższych częstotliwościach. A wyższa częstotliwość przekłada się wprost na mniejszy dławik i mniejsze kondensatory – czyli na mniejszą, lżejszą i gęstsza energetycznie przetwornicę.

Do tego dochodzi odporność termiczna: szeroka przerwa energetyczna pozwala przyrządom WBG pracować przy wyższych temperaturach złącza niż krzem. Suma tych właściwości – niższe straty, wyższe częstotliwości, większa gęstość mocy, lepsza praca w wysokiej temperaturze – sprawia, że materiały szerokoprzerwowe nie są kosmetycznym ulepszeniem krzemu, lecz zmianą jakościową. Inżynierowie opisują tę przewagę za pomocą tak zwanych współczynników jakości (*figures of merit*), jak współczynnik Baliga czy Johnsona; dla SiC i GaN przyjmują one wartości o rząd wielkości korzystniejsze niż dla krzemu.

Nie znaczy to, że krzem odchodzi. Pozostaje tańszy, ma w pełni dojrzały łańcuch dostaw i w wielu zastosowaniach niskonapięciowych oraz niskoczęstotliwościowych jest w zupełności wystarczający. Technologie WBG opłacają się tam, gdzie sprawność, gęstość mocy albo praca przy wysokim napięciu mają realną wartość – i to właśnie ten obszar rynku rośnie dziś najszybciej.

Kluczowa dla tego artykułu jest jednak inna obserwacja. Półprzewodniki szerokoprzerwowe to nie jeden materiał, lecz dwa – a SiC i GaN, choć dzielą zaletę szerokiej przerwy energetycznej, znacząco różnią się między sobą. Węglik krzemu dojrzał najpierw w obszarze wysokonapięciowym:

w motoryzacji, trakcji i przemysłowych zasilaczach dużej mocy. Azotek galu wszedł na rynek od strony niskich napięć i wysokich częstotliwości – w ładowarkach, zasilaczach i przetwornicach o dużej gęstości mocy. Przez pewien czas technologie te niemal ze sobą nie konkurowały. Dziś, gdy oba materiały są dostępne w klasie 650 V, wreszcie się spotkały – i cała dalsza część artykułu poświęcona jest właśnie temu spotkaniu.

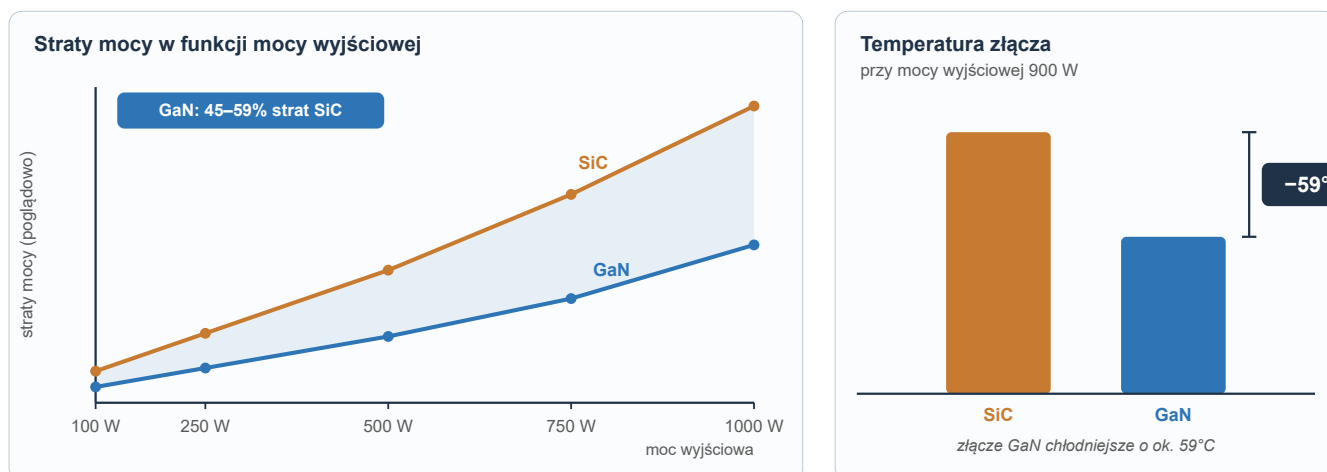
Dlaczego akurat 200 W to pole bitwy

W literaturze przedmiotu od dawna funkcjonuje wygodny podział: SiC to następca krzemu w domenie wysokonapięciowej – powyżej 1200 V – a GaN w domenie niskonapięciowej, poniżej 650 V. Tak właśnie ujmują rzecz autorzy porównawczych badań konwerterów, traktując MOSFET SiC i tranzystor GaN HEMT jako sukcesorów krzemu odpowiednio w obszarze średnich-wysokich i niskich napięć. Kłopot w tym, że wraz z pojawieniem się przyrządów obu typów o napięciu 650 V te dwie domeny zaczęły się pokrywać – i to dokładnie tam, gdzie pracuje wiele praktycznych przetwornic: zasilacze serwerowe, sekcje DC/DC, ładowarki, układy przemysłowe.

Klasa 200 W przy napięciu wejściowym nieprzekraczającym 650 V jest modelowym przykładem tej strefy spornej. Nie jest to już domena, w której GaN nie ma czego szukać – jak w trakcyjnym inwerterze pojazdu elektrycznego przy 800...1000 V – ani taka, w której SiC pozostaje bezkonkurencyjny. Obie technologie wchodzą tu z pełnym portfolio argumentów, a to oznacza, że decyzja wymaga zrozumienia nie sloganów marketingowych, lecz fizyki przełączania.

Straty i temperatura złącza: GaN obok SiC

Ta sama synchroniczna przetwornica buck 400 V/200 V, 200 kHz — porównanie poglądowe



Rysunek 3. Straty mocy i temperatura złącza układu z GaN i z SiC w tej samej synchronicznej przetwornicy buck – zestawienie poglądowe na podstawie danych porównawczych

200 kHz najwyższą sprawność – rzędu 91...94% – osiągnął wariant z azotkiem galu.

Warto przy tym pamiętać, skąd te różnice się biorą. W przetwornicy o umiarkowanej mocy i niewysokim napięciu większość strat tranzystora to straty przełączania – energia tracona w krótkich chwilach, gdy przyrząd przechodzi między stanem zatkania a przewodzeniem. To właśnie w tych chwilach przewaga GaN – mały ładunek bramki, brak ładunku odzysku, niewielkie pojemności – przekłada się na realne waty. Im wyższa częstotliwość, tym więcej takich przejść w jednostce czasu i tym wyraźniej rozjeżdżają się bilanse obu technologii.

Te liczby nie oznaczają jednak, że SiC „przegrywa”. Oznaczają coś bardziej precyzyjnego: w twardym przełączaniu, przy umiarkowanych prądach i napięciu nieprzekraczającym 650 V, GaN ma fizyczną przewagę. Gdy jednak warunki się zmieniają – rośnie napięcie, prąd albo temperatura otoczenia – przewaga ta topnieje, by w końcu zniknąć. Dlatego zamiast pytać „co jest lepsze”, warto zapytać „dlaczego”, i przyjrzeć się parametrom, które o wszystkim decydują.

Cztery parametry, które decydują

Skąd bierze się przewaga GaN i kiedy się ona kończy? Decyduje o tym kilka parametrów, które warto rozważyć po kolei.

Ładunek odzysku diody (Qrr). To najmocniejszy pojedynczy argument GaN. Tranzystor GaN HEMT nie ma ładunku gromadzonego w pasożytniczej diodzie – jego Qrr jest praktycznie zerowy. W topologiach z twardym przełączaniem (mostek, synchroniczny buck) eliminuje to całą klasę strat i prądów udarowych, które w MOSFET-cie SiC trzeba uwzględnić w bilansie. W przetwornicy 200 W z mostkiem bywa to czynnik przesądzający.

Ładunek bramki (Qg) i pojemność wyjściowa (Coss).

GaN ma typowo niższy ładunek bramki i mniejsze pojemności pasożytnicze, co oznacza szybsze przełączanie i niższe straty sterowania. To bezpośrednie źródło przewagi przy wysokich częstotliwościach – i dzięki temu przetwornicę GaN można „podkręcić” tak, by radykalnie zmniejszyć rozmiar dławika i kondensatorów filtrujących.

Rezystancja w stanie włączenia w funkcji temperatury (Rds(on) vs T). Tu obraz się komplikuje. Rds(on) obu technologii rośnie z temperaturą, ale charakterystyki różnią się producent od producenta. Noty aplikacyjne wytwórców, którzy mają w portfolio jednocześnie SiC i GaN, pokazują tę zależność wprost – i to ona, a nie wartość katalogowa w 25°C, decyduje o rzeczywistych stratach przewodzenia w nagrzanym układzie.

Przewodność cieplna materiału. To jedyny z czterech parametrów, w którym SiC ma wyraźną, fizyczną przewagę. Wynika ona z samego materiału – i wróćmy do niej w sekcji poświęconej termice.

REKLAMA



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

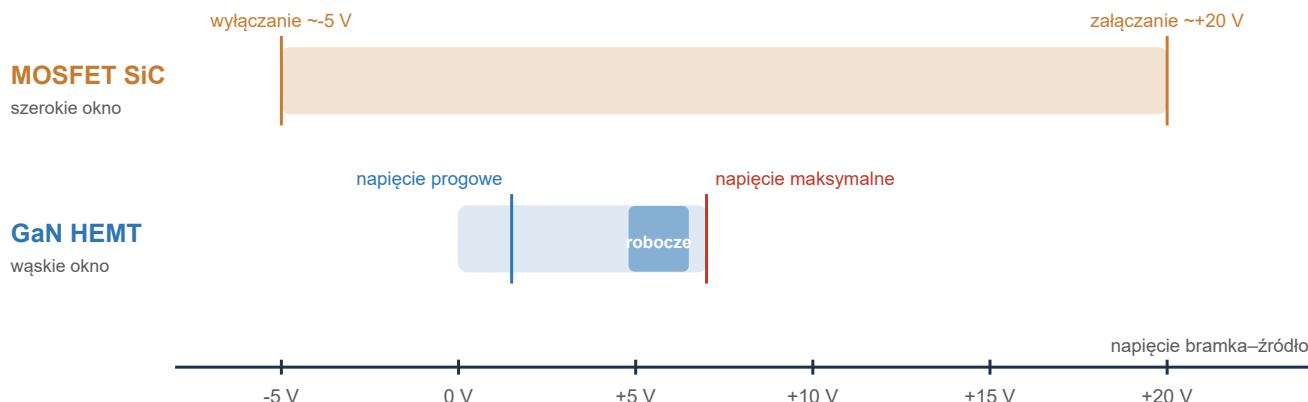


✉ bornico@bornico.com.pl 🌐 www.bornico.com.pl

☎ +48 517 312 709 | +48 517 312 419

Okno bezpiecznego sterowania bramką

Dlaczego GaN wybaczta mniej niż SiC — zakresy napięcia bramki



Zakresy poglądowe; dokładne poziomy zależą od rodziny przyrządów i karty katalogowej. Pomyłka o 1–2 V, nieszkodliwa dla SiC, w GaN mieści się już poza oknem.

Rysunek 4. Okno bezpiecznego napięcia bramki – szerokie dla MOSFET-a SiC, wąskie dla tranzystora GaN

Te cztery parametry nie działają w oderwaniu od siebie. Trzypierwsze pracują na korzyść GaN, czwarty – na korzyść SiC, a ich wypadkowa zależy od konkretnego punktu pracy. Dlatego nie istnieje uniwersalna odpowiedź; istnieje natomiast metoda doboru – i to do niej wrócimy w sekcji ze wskazówkami praktyków.

Sterowanie bramką – tu projekt wygrywa albo przegrywa

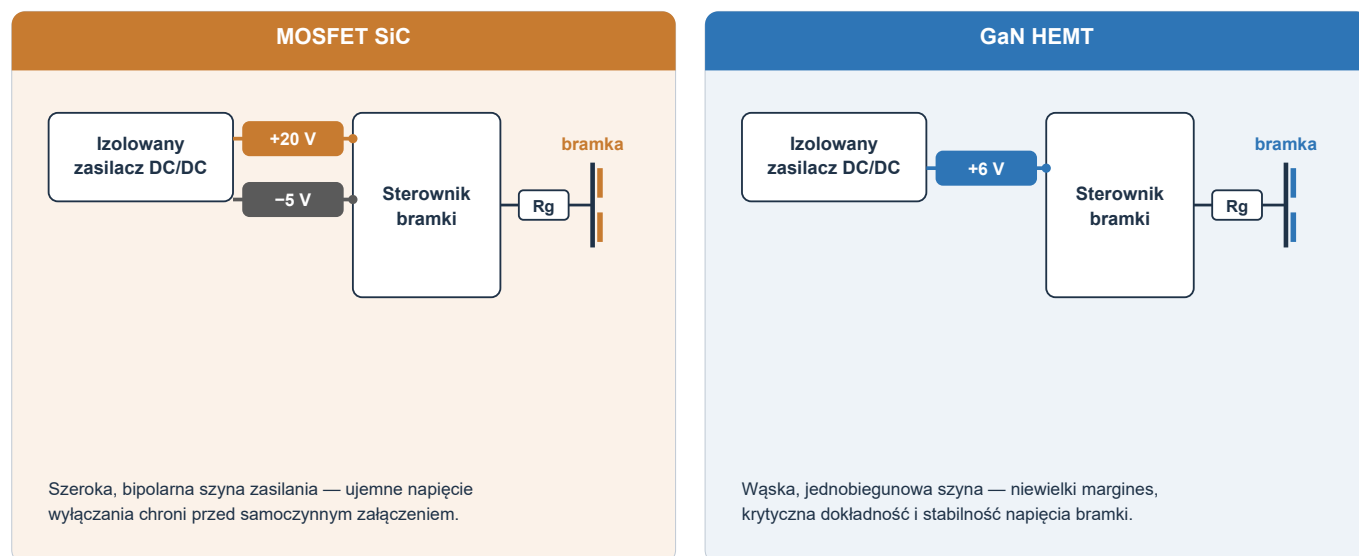
Jeśli ten artykuł ma przekazać jedną praktyczną przestrożę, brzmi ona tak: o powodzeniu przetwornicy z przyrządami szerokoprzerwowymi nie decyduje wybór tranzystora, lecz projekt obwodu jego bramki. Eksperci od zasilania sterowników bramkowych zwracają uwagę, że obie technologie mają tu zupełnie różne wymagania.

MOSFET SiC potrzebuje stosunkowo wysokiego napięcia bramki – typowo +20 V do włączania oraz napięcia ujemnego rzędu –5 V do pewnego wylączenia, co chroni przed samoczynnym załączeniem (parasitic turn-on). Tranzystor GaN HEMT typu e-mode pracuje przy napięciu znacznie niższym, rzędu +6 V, i – co kluczowe – ma bardzo wąskie okno bezpieczne między napięciem progowym a maksymalnym napięciem bramki. Pomyłka o 1...2 V, która MOSFET-owi SiC nie zaszkodzi, tranzystor GaN potrafi zniszczyć.

Z tego wynikają konkretne konsekwencje projektowe: inny sterownik, inny izolowany zasilacz sterownika, inne marginesy bezpieczeństwa. Materiały producentów sterowników podkreślają, że zasilacz bramki dla SiC i dla GaN to dwa różne zadania projektowe – od poziomów napięć, przez wymagania izolacji, po odporność na strome zbocza.

Zasilanie sterownika bramkowego: SiC i GaN

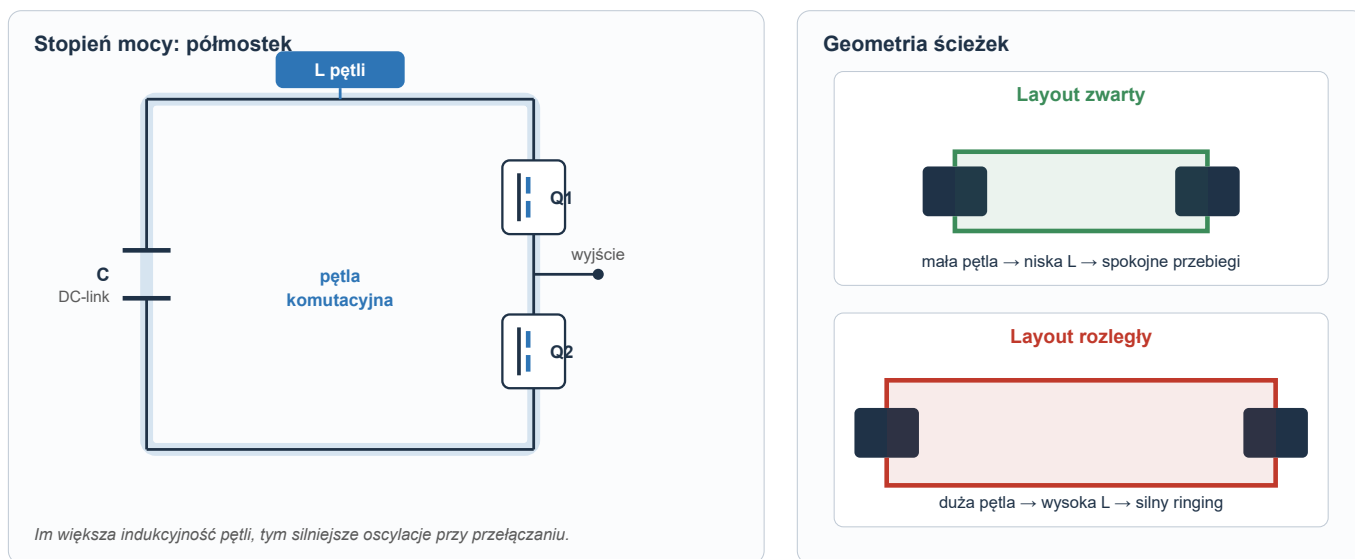
Dwa różne zadania projektowe — poziomy napięć i architektura szyny bramki



Rysunek 5. Zasilanie sterownika bramkowego – bipolarna szyna SiC (+20 V/–5 V) wobec wąskiej, jednobiegowej szyny GaN (ok. +6 V)

Pętla komutacyjna — dlaczego layout decyduje

Indukcyjność pasożytnicza pętli rezonuje z pojemnościami i generuje oscylacje



Rysunek 6. Pętla komutacyjna półmostka i wpływ geometrii ścieżek – zwarty layout ogranicza indukcyjność pasożytniczą i oscylacje

Osobnym, dobrze udokumentowanym problemem są oscylacje na bramce. Przy bardzo szybkich zboczach indukcyjność pętli bramki rezonuje z pojemnościami tranzystora. SiC bywa pod tym względem „wybaczący”, a GaN – wymagający: prace przeglądowe poświęcone strategiom poprawy parametrów dyskretnych przyrządów WBG wskazują oscylacje bramki GaN jako jeden z głównych problemów projektowych, który trzeba opanować doborem rezystora bramki, koralikiem ferrytowym i – przede wszystkim – geometrią ścieżek.

Pasożyty, ringing (dzwonienie) i sztuka layoutu

Przewaga GaN to przewaga szybkości – a szybkość ma swoją cenę. Przy zboczach krótszych niż nanosekunda nawet kilka nanohenrów indukcyjności pasożytniczej staje się problemem. Analizy przełączania sub-nanosekundowego pokazują skalę zjawiska: indukcyjności pasożytnicze rzędu 1...5 nH w pętli komutacyjnej rezonują z pojemnością wyjściową rzędu 10...100 pF, generując oscylacje o częstotliwościach sięgających setek megaherców.

Dla konstruktora przetwornicy 200 W oznacza to jedno: jeśli wybrano GaN, projekt obwodu drukowanego prze staje być czynnością rutynową, a staje się integralną częścią projektu elektrycznego. Pętla komutacyjna musi być minimalna, kondensatory wejściowe – umieszczone tuż przy tranzystorach, a obwód bramki – możliwie najkrótszy. SiC, przełączający wolniej, jest pod tym względem bardziej tolerancyjny – i bywa to realnym argumentem za jego wyborem, jeśli zespół nie ma doświadczenia z layoutem szybkich przetwornic albo płytka musi powstać szybko i tanio.

To jedna z tych prawd, których „nie mówią na studiach”: karta katalogowa obiecuje sprawność, ale dostarcza ją dopiero dobrze zaprojektowana płytka.

Termika – cichy atut węgla krzemu

We wszystkich dotychczasowych sekcjach przewagę najczęściej zyskiwał GaN. Termika tę tendencję odwraca. Raport amerykańskiego Oak Ridge National Laboratory, porównujący materiały szerokoprzerwowe pod kątem współczynników jakości (figures of merit), pokazuje wprost, że pod względem przewodności cieplnej GaN wypada gorzej niż SiC. Nie jest to kwestia konkretnego producenta ani generacji przyrządu – to właściwość samego materiału.

Konsekwencja jest czysto praktyczna. SiC łatwiej odprowadza ciepło z obszaru złącza, lepiej znosi wysokie temperatury otoczenia i pracę w pobliżu granic obciążenia. Dlatego przewodniki doboru publikowane przez

REKLAMA

PRODUCENT
ELEMENTÓW
INDUKCYJNYCH

www.feryster.pl

FERYSTER

Tabela 1. GaN HEMT a MOSFET SiC w klasie 650 V – zestawienie cech istotnych dla przetwornicy 200 W

Parametr/cecha	GaN HEMT 650 V	MOSFET SiC 650 V
Ładunek odzysku Qrr	praktycznie zerowy	niezerowy – trzeba uwzględnić
Ładunek bramki Qg	niższy	wyższy
Napięcie sterowania bramki	ok. +6 V, wąskie okno	ok. +20 V/-5 V
Pojemności pasożytnicze Coss	mniejsze	większe
Przewodność cieplna materiału	niższa	wyższa
Typowa częstotliwość pracy	bardzo wysoka	wysoka
Odporność termiczna	dobra	bardzo dobra
Tolerancja layoutu PCB	wymagająca	bardziej wybacząca
Kwalifikacje motoryzacyjne	rosnące	ugruntowane
Atut w klasie 200 W	sprawność, gęstość mocy	termika, odporność, prostota
Zestawienie ma charakter poglądowy – konkretne wartości należy odczytać z kart katalogowych wybranych przyrządów.		

producentów SiC rekomendują tę technologię wszędzie tam, gdzie liczy się odporność termiczna, wysoka temperatura pracy oraz napięcia od 650 V wzwyż, a GaN – tam, gdzie priorytetem jest wysoka częstotliwość przy napięciu rzędu 600 V.

Różnica ta nie przekreśla GaN – w wielu przetwornicach 200 W pracujących w umiarkowanych warunkach zapas termiczny azotku galu jest w zupełności wystarczający. Staje się jednak istotna na granicy obciążenia: tam, gdzie konstruktor projektuje układ blisko jego limitów cieplnych, lepsze odprowadzanie ciepła przez SiC daje margines, który trudno kupić w inny sposób.

Dla przetwornicy 200 W pytanie brzmi więc bardzo konkretnie: czy układ będzie pracował w klimatyzowanej szafie, czy w gorącej, zamkniętej obudowie bez wymuszonego chłodzenia? W tym drugim przypadku

fizyka materiału zaczyna przemawiać za SiC – nawet jeśli sprawność zmierzona w temperaturze pokojowej była nieco niższa.

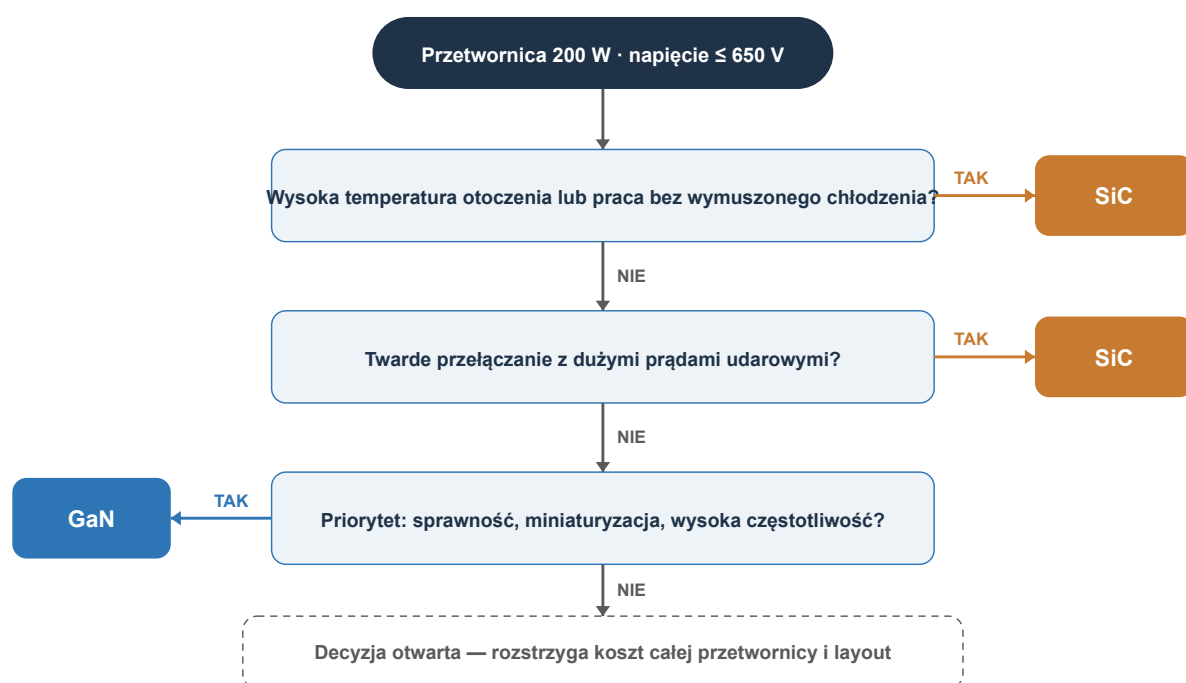
Wskazówki praktyków: jak podjąć decyzję

Zbierzmy teraz w jedną całość to, co o wyborze technologii radzą praktycy.

Najpierw zdefiniuj warunki brzegowe, nie technologii. Producenci mający w portfolio jednocześnie SiC i GaN – a więc pozbawieni interesu w promowaniu jednej tylko strony – formułują to spójnie: o wyborze decyduje aplikacja, a nie materiał. Zaczniij od czterech liczb: napięcie wejściowe, częstotliwość przełączania, maksymalna temperatura otoczenia oraz topologia (twarde czy miękkie przełączanie).

Ścieżka decyzji dla przetwornicy 200 W

Kolejność pytań, którą podpowiadają praktycy — od warunków brzegowych do technologii



Rysunek 7. Ścieżka decyzji: kolejność pytań prowadząca od warunków brzegowych do wyboru technologii

Dopasuj technologię do topologii. Jeśli projekt zakłada miękkie przełączanie – rezonansowe topologie LLC lub CLLC – atutem stają się małe pojemności i niski ładunek bramki, a wymagania wobec layoutu łagodnieją; tu GaN czuje się znakomicie i jest często wprost rekomendowany. Jeśli natomiast topologia jest twardo przełączana i prądowo wymagająca, do gry wracają odporność i termika SiC.

Policz koszt całego rozwiązania, nie samego tranzystora. GaN bywa droższy w przeliczeniu na sztukę, ale eliminuje elementy – mniejszy dławik, mniej chłodzenia, czasem prostszą topologię jednostopniową. SiC bywa tańszy jako komponent, lecz może wymagać solidniejszego radiatora. Liczy się rachunek na poziomie kompletnej przetwornicy.

Zasymuluj, zanim zbudujesz. Klasyczna wskazówka ekspertów od impulsowych zasilaczy – przypomnijmy choćby dorobek Christophe’a Basso, autora podręcznika o symulacjach SPICE w projektowaniu SMPS – brzmi: przetwornicę z przyrządami WBG warto przepuścić przez symulator, zanim powstanie pierwsza płytką. Strone zbocza GaN i rezonanse pasożytnicze są znacznie tańsze do wykrycia w modelu niż na stanowisku pomiarowym.

Nie lekceważ łańcucha dostaw i kwalifikacji. SiC ma za sobą dłuższą historię w aplikacjach motoryzacyjnych i przemysłowych. Jeśli przetwornica 200 W ma trafić do wyrobu o długim cyklu życia, dostępność drugiego źródła i dojrzałość kwalifikacji bywają argumentem nie mniej ważnym niż ułamek procenta sprawności.

Werdykt dla przetwornicy 200 W

Wróćmy do pytania z tytułu. Z całości materiału – debat, not aplikacyjnych, pomiarów i raportów – wyłania się odpowiedź zaskakująco spójna, choć nie jednoznaczna.

Przy 200 W i napięciu nieprzekraczającym 650 V, w typowej aplikacji o umiarkowanej temperaturze otoczenia,

Reguła kciuka – szybka ścieżka decyzji

- Napięcie ≤ 650 V, wysoka częstotliwość, miniaturyzacja klucza $\rightarrow$ GaN.
- Wysoka temperatura otoczenia lub praca bez wymuszonego chłodzenia $\rightarrow$ SiC.
- Twarde przełączanie, duże prądy udarowe, potrzebny zapas wytrzymałości $\rightarrow$ SiC.
- Topologia rezonansowa LLC/CLLC, priorytetem maksymalna sprawność $\rightarrow$ GaN.
- Brak doświadczenia z layoutem szybkich przetwornic, presja czasu $\rightarrow$ SiC.

Zawsze: policz koszt kompletnej przetwornicy, a nie samego tranzystora – i zasymuluj projekt, zanim powstanie pierwsza płytką.



z twardym lub miękkim przełączaniem przy podwyższonej częstotliwości – GaN jest zwykle lepszym wyborem. Daje wyższą sprawność, niższą temperaturę złącza i większą gęstość mocy, a zerowy Qrr upraszcza topologie mostkowe. Ceną są wymagający layout i ostrożne sterowanie bramką.

SiC staje się lepszym wyborem, gdy choć jeden z warunków brzegowych się zaostrza: wysoka temperatura otoczenia lub praca bez wymuszonego chłodzenia, napięcia powyżej 900 V – na przykład w szerszej rodzinie wyrobów – duże prądy w twardym przełączaniu, albo gdy priorytetem są prostota projektu i szybkie, tolerancyjne uruchomienie.

Granica decyzyjna leży przy 650 V – i, jak pokazał niemal remisowy werdykt sali APEC, pozostaje przedmiotem żywej dyskusji. Dla konstruktora płynie z tego wniosek najcenniejszy: w klasie 200 W nie ma wyboru „domyślnego”. Jest tylko świadoma decyzja, podjęta po zważeniu czterech liczb – napięcia, częstotliwości, temperatury i topologii – oraz po uczciwym rachunku kosztu całego rozwiązania.

WM

Bibliografia – źródła i materiały pogłębiające

- [1] M. Di Paolo Emilio, A. Shaukat, „The Great Debate at APEC 2025: GaN vs. SiC”, Power Electronics News, 2025.
- [2] „WAWT’s Debrief on APEC 2025”, Wired & Wireless Technologies, kwiecień 2025.
- [3] Materiały programowe 40. konferencji IEEE Applied Power Electronics Conference (APEC 2025), Atlanta, sesje debat (rap sessions).
- [4] J. Xu, „A Performance Comparison of GaN E-HEMTs versus SiC MOSFETs in Power Switching Applications”, GaN Systems/eeper.com.
- [5] „Comparison of Si, SiC, and GaN based DC-DC converters”, Journal of Electrical Engineering, 2024.
- [6] „Performance Evaluation of GaN-Based Synchronous Boost Converter under Various Output Voltage, Load Current and Switching Frequency Operations”, publikacja naukowa.
- [7] Texas Instruments, „Performance and benefits of GaN versus SiC”, nota aplikacyjna SLYT801.
- [8] Wolfspeed, „SiC vs. GaN Selection Guide”, dokument techniczny.
- [9] Infineon Technologies, materiały o rodzinie CoolSiC (przyszyty SiC).
- [10] Infineon Technologies, materiały o rodzinie CoolGaN (przyszyty GaN).
- [11] RECOM Power, „Gate Driver Power Supply Challenges for SiC & GaN Wide-Bandgap Devices”.
- [12] Oak Ridge National Laboratory, „Comparison of Wide-Bandgap Semiconductors for Power Electronics Applications”.
- [13] „Review and Outlook on GaN and SiC Power Devices”, IEEE Journal of Emerging and Selected Topics in Power Electronics.
- [14] „Performance Improvement Strategies for Discrete WBG Power Devices”, Frontiers in Energy Research, 2021.
- [15] Ch. P. Basso, „Switch-Mode Power Supplies: SPICE Simulations and Practical Designs”, McGraw-Hill.
- [16] Materiały techniczne EPC (Efficient Power Conversion) dotyczące tranzystorów GaN typu e-mode.
- [17] Analiza przełączania sub-nanosekundowego przyrządów GaN – wpływ ich materiałowych i procesów wytwarzania GaN oraz SiC.
- [19] Materiały producentów sterowników bramkowych dotyczące wymagań sterowania przyrządów SiC i GaN.
- [20] Branżowe omówienia konferencji APEC 2025 – trendy w konwersji mocy dla centrów danych.

Analiza, symulacja i eksperyment – klucz do udanego projektu kompensacji zasilacza

Sterowanie pętlą regulacji to istotny aspekt projektowania przetwornicy impulsowej. Z różnych powodów jego analizę często odkłada się jednak na sam koniec projektu, gdy główne podzespoły zostały już dobrane. Posługując się zwykłą metodą prób i błędów, można niekiedy odnieść wrażenie, że konstrukcja zapewniająca akceptowalną odpowiedź przejściową na oscyloskopie jest gotowa do produkcji – jest to jednak postawa bardzo nierozważna i potencjalnie kosztowna.

Dzieje się tak dlatego, że większość elementów użytych w przetwornicy obciążona jest pasożytniczymi składowymi, których rozległe skutki pozostają ukryte na etapie prototypu. Bez gruntownej analizy popartej symulacjami i pomiarami pętli nie mamy pojęcia, jak wyglądają zapasy fazy i wzmocnienia ani jak są one stabilne. Bardzo prawdopodobne, że tak niedbale zaprojektowana przetwornica zawiedzie w produkcji lub wkrótce po uruchomieniu w terenie. Aby uchronić Cię przed takimi sytuacjami, w artykule omówiono niektóre z dostępnych obecnie narzędzi pozwalających obliczyć, zasymulować i zmierzyć pętlę regulacji prototypu, zanim bezpiecznie naciśniesz przycisk rozpoczęcia produkcji.

Znaczenie odpowiedzi stopnia mocy

Przetwornica impulsowa składa się ze stopnia mocy, którego wyjście jest sterowane zmienną napięciową. Oznaczona w tym artykule jako V_{err} lub V_c , jest ona dostarczana przez blok kompensacji utrzymujący wyjście przetwornicy w zakresie regulacji. W przypadku przetwornic pracujących ze stałą częstotliwością przełączania F_{sw} zmienną sterującą jest współczynnik wypełnienia D . Nie zawsze jednak tak jest – niektóre przetwornice są sterowane zmienną częstotliwością (np. przetwornice rezonansowe typu LLC) albo zmiennym czasem załączenia lub wyłączenia. W artykule skupimy się głównie na układach o stałej częstotliwości przełączania.

Napięcie błędu V_{err} może sterować współczynnikiem wypełnienia bezpośrednio – mówimy wtedy o sterowaniu w trybie napięciowym (VM, voltage-mode) lub



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.

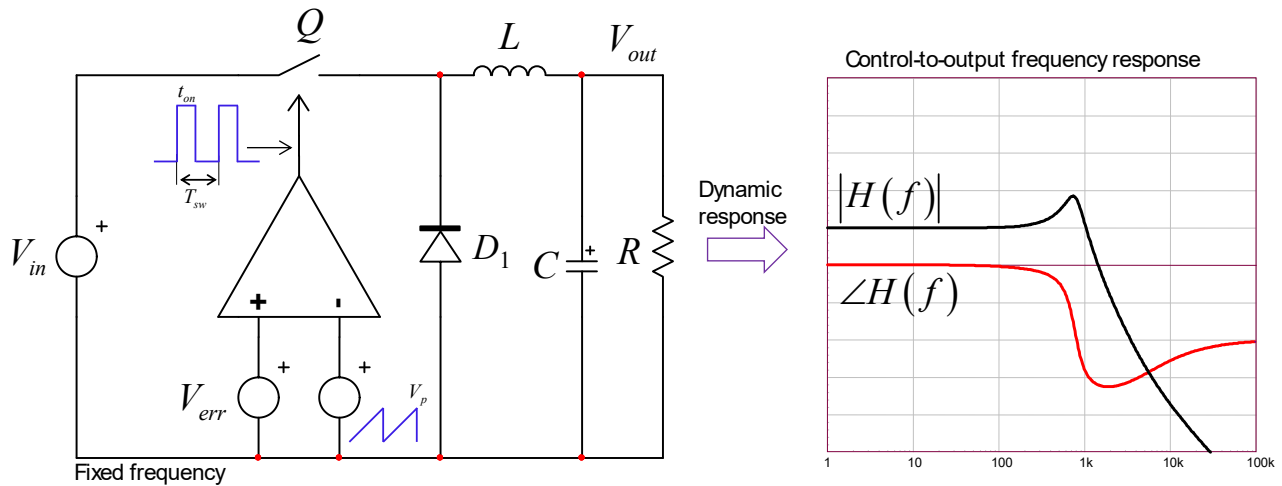
Basso jest postacią niezwykle cenioną w globalnym środowisku inżynierskim. Znany jest jako autor jedenastu książek technicznych, w tym popularnych publikacji na temat metod szybkiej analizy obwodów oraz symulacji układów zasilających. Jego wkład w branżę to nie tylko teoria, ale przede wszystkim praktyczne 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.

o bezpośrednim sterowaniu wypełnieniem. Z drugiej strony istnieje również sterowanie w trybie prądowym (CM, current-mode), w którym napięcie sterujące V_c ustala szczytowy prąd cewki cykl po cyklu za pośrednictwem rezystora pomiarowego, a tym samym pośrednio wyznacza roboczy współczynnik wypełnienia.

Jednak gdy na oscyloskopie wyświetlane są przebiegi z przetwornicy, nie da się rozpoznać, czy pracuje ona w trybie CM, czy VM. Wynika to z tego, że stopień mocy jest w obu strukturach podobny – zmienia się jedynie sposób wyznaczania współczynnika wypełnienia: przetwornica buck dostarczająca 5 V ze źródła 10 V do obciążenia wykaże teoretyczne wypełnienie 50% niezależnie od tego, czy układ pracuje w trybie napięciowym, czy prądowym.

Naszym celem jako projektantów zasilaczy jest zbudowanie niezawodnej przetwornicy zdolnej dostarczać precyzyjnie regulowane napięcie (lub prąd), niewrażliwej zarazem na warunki pracy: wahania napięcia źródła, zmiany temperatury otoczenia, różne stany obciążenia itd. Poza tymi praktycznymi wymaganiami projektant musi zapewnić, by jego przetwornica pozostawała stabilna i sprawna przez cały okres eksploatacji.

Trzeba liczyć się z naturalnym rozrzutem produkcyjnym oraz degradacją parametrów elementów wskutek



Rysunek 1. Poszukujemy dynamicznej odpowiedzi stopnia mocy

starzenia. Co stanie się z moimi pięknymi zapasami projektowymi za pięć lat? Na ile pewny jestem swoich wyborów, gdy zaopatrzeniowiec pokaże mi nowy, tańszy kondensator wybrany przez fabrykę? „Hej Joe, możesz potwierdzić, że nowa milionowa partia Twoich zasilaczy będzie w porządku, jeśli na kondensator wyjściowy wybierzemy markę B zamiast obecnie stosowanej marki A?”. Czy potrafiłbyś bez wahania odpowiedzieć na to pytanie?

Owszem, potrafiłbyś – jeśli odrobiłeś pracę domową i starannie przeanalizowałeś wpływ pasożytów kondensatora na częstotliwość zerową przejścia (crossover) oraz na zapas fazy. Ale jeśli tego nie zrobiłeś i tylko obserwowałeś odpowiedź skokową na stanowisku, kręcąc gałkami R i C kompensatora – to cóż, wytrzymaj kroplę potu z czoła, bo czekają Cię krótkie noce, byle tylko uniknąć katastrofy.

Jednym ze sposobów uniknięcia pułapki jest postępowanie zgodne ze sztuką i rozpoczęcie od odpowiedzi stopnia mocy. To jedyny właściwy punkt wyjścia: zanim w ogóle pomyślisz o strategii sterowania, musisz scharakteryzować układ, którym chcesz sterować. Chodzi o ustalenie, jak zmienna wyjściowa reaguje na zmianę sygnału sterującego. Oznacza to, że potrzebujesz transmitancji sterowanie-wyjście (control-to-output) projektowanej przetwornicy buck lub boost: jaka jest dynamiczna odpowiedź V_{out} na zadane pobudzenie w V_{err} (rysunek 1). Innymi słowy – jaka jest odpowiedź obiektu (plant)?

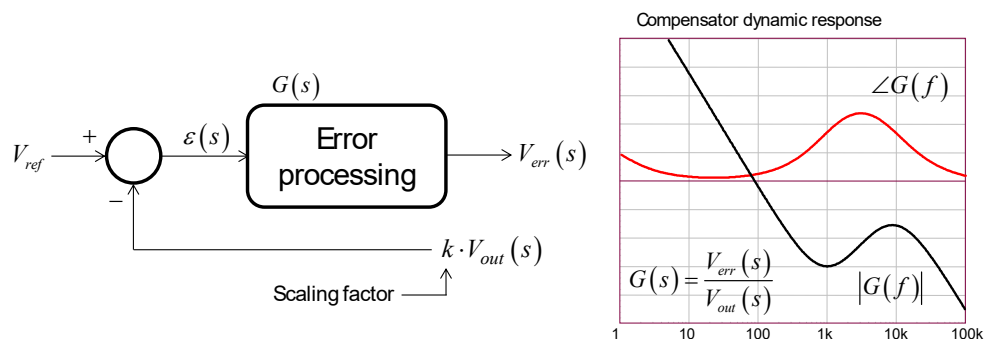
Gdy masz już w ręku wykres modułu i fazy transmitancji, możesz pomyśleć o strategii kompensacji polegającej na umieszczeniu biegunów, zer oraz wzmocnienia (lub tłumienia) w wybranych częstotliwościach tak, aby spełnić założenia projektowe, w tym odpowiednie zapasy fazy i wzmocnienia. Rolę

kompensatora w torze sprzężenia zwrotnego zasilacza ilustruje rysunek 2 wraz z przykładową odpowiedzią kompensatora.

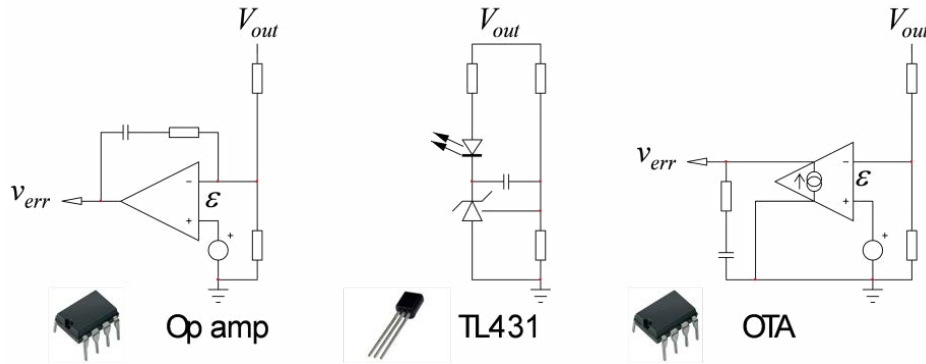
Budując kompensatory, można pójść kilkoma drogami, co ilustruje rysunek 3. Klasyczne podejście, szeroko opisane w literaturze, wykorzystuje wzmacniacz operacyjny do zbudowania filtra – kompensator jest bowiem niczym innym jak filtrem aktywnym.

W przemyśle króluje jednak regulator równoległy TL431, który można znaleźć w zdecydowanej większości sprzedawanych dziś zasilaczy sieciowych. Przyznaję, że trudno go pobić pod względem prostoty i kosztu: za kilka centów otrzymujesz wzmacniacz operacyjny o umiarkowanie dużym wzmocnieniu w pętli otwartej (55 dB), wyposażony w precyzyjne źródło napięcia odniesienia V_{ref} 2,5 V (a w wersji TLV nawet 1,24 V). Element dostępny jest w różnych obudowach, a niektóre wersje akceptują napięcie do 36 V. Dobór tego układu wiąże się jednak z innymi zagadnieniami związanymi z torem szybkim i wolnym, opisanymi w pozycji [1].

Do celów kompensacji można wybrać również wzmacniacz transkonduktancyjny (OTA). Projektanci układów scalonych lubią OTA, ponieważ zajmują one mniejszą powierzchnię krzemu niż odpowiedniki na wzmacniaczach operacyjnych. Osobiście nie jestem wielkim



Rysunek 2. Kompensator (czyli blok przetwarzania błędów oraz współczynnik skalujący) to miejsce, w którym wstawia się bieguny i zera oraz kształtuje pożądaną odpowiedź częstotliwościową



Rysunek 3. Projektując kompensator, można wybierać spośród kilku elementów aktywnych

zwoleńnikiem OTA, ponieważ nie mamy w nim masy pozorowanej (virtual ground), jaką zapewnia kompensator oparty na wzmacniaczu operacyjnym. Co więcej, stosunek dzielnika rezystancyjnego wpływa na rozmieszczenie biegunów i zer.

OTA są popularne w układach korekcji współczynnika mocy (PFC) i dobrze nadają się do realizacji kompensatorów o umiarkowanym podbiciu fazy. Jeśli jednak planujesz je zastosować tam, gdzie wymagane jest duże podbicie fazy, możesz natrafić na górne ograniczenie narzucone przez stosunek V_{out}/V_{ref} opisane w pozycji [2].

Podbicie fazy (phase boost) to ilość dodatkowej fazy, jakiej potrzebujesz od kompensatora, aby spełnić założenia dotyczące zapasu fazy – zwykle wartość większa niż 45° . Na rysunku 4 widać stopień mocy o opóźnieniach fazowych 90° lub 145° przy wybranych częstotliwościach f_1 i f_2 . Gdybyś zamknął pętlę zwykłym integratorem o stałym opóźnieniu 270° , to suma obu składników przy f_1 wyniesie -360° (czyli 0°): sygnał wraca w fazie do punktu wstrzyknięcia i spełnione są warunki podtrzymania oscylacji. Nie jest to coś, czego byś chciał – chyba że celem jest zbudowanie generatora.

Gdybyś teraz wymusił przejście przez 0 dB (crossover) przy f_2 , zapas fazy będzie ujemny, co oznacza, że bieguny

w pętli zamkniętej leżą w prawej półpłaszczyźnie – układ jest niestabilny. Problem ten zwalcza się, tworząc podbicie fazy przy f_1 lub f_2 . Wstawiając bieguny i zera w kompensatorze, kształtujesz jego odpowiedź fazową tak, by nie była już stale równa -270° , lecz mniejsza. Po dodaniu jej do odpowiedzi obiektu łączna faza będzie większa (mniej ujemna) niż -360° , co tworzy pożądany zapas fazy zapewniający stabilność.

Możemy wyróżnić trzy typy kompensatorów, znane jako typ 1, 2 i 3 [3]. Przedstawiono je na rysunku 5.

Pierwszy typ zawiera biegun w początku układu współrzędnych: jest to integrator opisany następującą transmitancją:

$$G_1(s) = -\frac{1}{\frac{s}{\omega_{po}}}$$

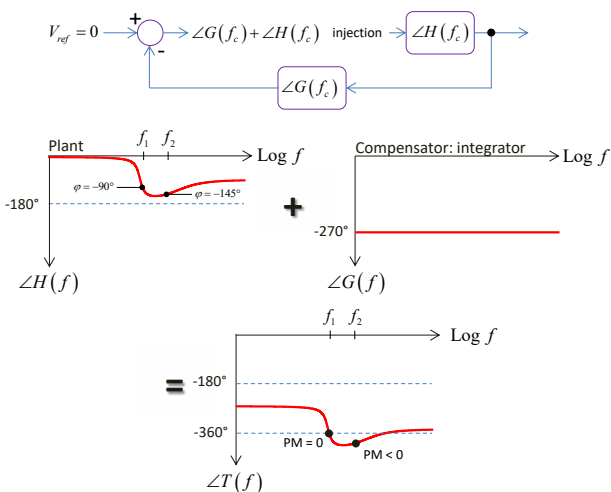
Nie ma tu podbicia fazy, a faza jest fazą odwracającą struktury wzmacniacza operacyjnego (-180°) powiększoną o fazę bieguna w zerze (-90°), co daje wynik końcowy -270° (czyli 90°).

Typ 2 spotyka się powszechnie w projektach sterowanych prądowo, gdzie potrzebne jest podbicie fazy poniżej 90° . Zawiera on biegun w zerze oraz dodatkowy biegun i zero. Biegun w zerze ($s = 0$) teoretycznie kasuje błąd statyczny (odchyłkę stałoprądową między wartością zadaną a uzyskiwaną po zamknięciu pętli). Biegun ten występuje w zdecydowanej większości kompensatorów, istnieją jednak techniki – jak tzw. kształtowanie rezystancyjne na wyjściu (output resistive shaping) – które celowo go pomijają, godząc się na niewielką odchyłkę [4].

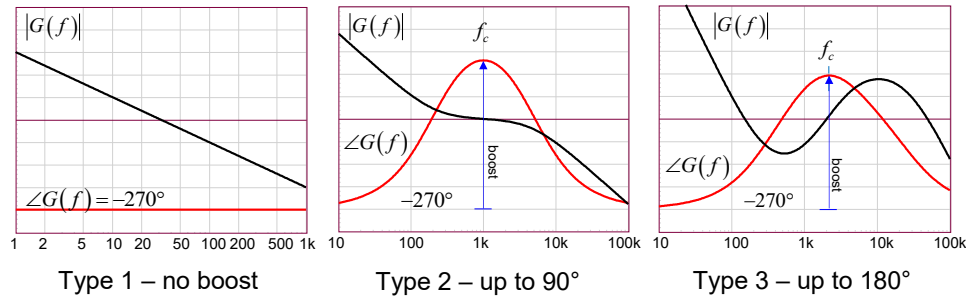
W typie 2 zero leży przed biegunem i podnosi fazę wraz ze wzrostem częstotliwości. Biegun włącza się później i podbicie fazy wraca do zera. Rozsuwając zero i biegun, regulujesz pożądane podbicie fazy aż do 90° . Zwróć uwagę, że zrównanie położenia bieguna i zera zamienia kompensator w typ 1 o zerowym podbiciu fazy.

Strukturę tę opisuje następująca transmitancja:

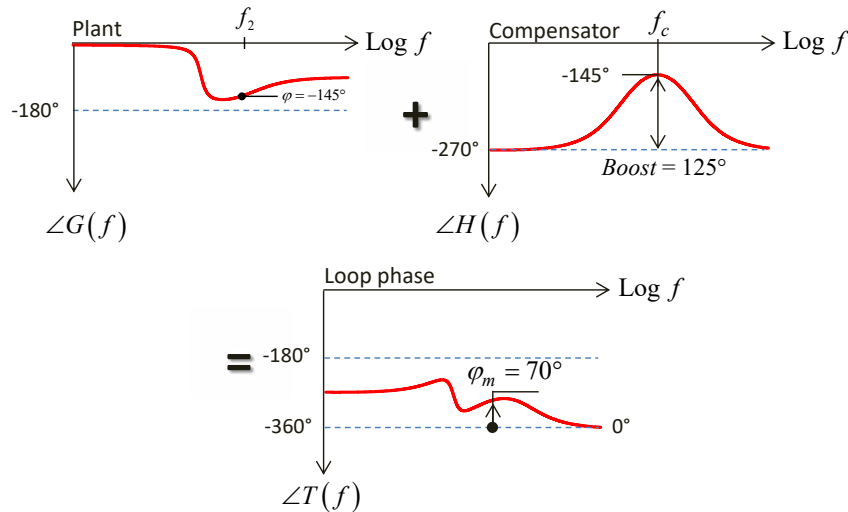
$$G_2(s) = -G_0 \frac{1 + \frac{\omega_z}{s}}{1 + \frac{s}{\omega_p}}$$



Rysunek 4. Zsumowanie fazy obiektu z fazą kompensatora powinno utrzymać łączne opóźnienie powyżej -360°



Rysunek 5. Trzy konfiguracje pozwalają zrealizować Twoją strategię kompensacji



Rysunek 6. Zapas fazy wynosi teraz 70° dzięki zastosowaniu kompensatora typu 3

W liczniku widać tzw. zero odwrócone (inverted zero), które pozwala wyłączyć przed nawias czynnik G_0 o wymiarze wzmocnienia [4].

Wreszcie kompensator typu 3 dodaje do typu 2 kolejną parę biegun-zero i potrafi podbić fazę aż do 180°. Opisuje go poniższe wyrażenie:

$$G_3(s) = -G_0 \frac{\left(1 + \frac{\omega_{z1}}{s}\right) \left(1 + \frac{s}{\omega_{z2}}\right)}{\left(1 + \frac{s}{\omega_{p1}}\right) \left(1 + \frac{s}{\omega_{p2}}\right)}$$

Jeśli w przykładzie z rysunku 4 zastosujemy teraz dla $G(s)$ układ typu 3 zamiast czystego integratora i podbijemy fazę o 125°, łączna faza pętli oddali się od granicy 0° (czyli -360°) i uzyskamy zapas 70°, jak pokazano na rysunku 6.

Można wyprowadzić wzór wiążący wymagane podbicie fazy z opóźnieniem fazowym stopnia mocy oraz pożądanym zapasem fazy φ_m . Wiemy, że odwracający wzmacniacz operacyjny oraz biegun w zerze wnoszą opóźnienie 270°, do którego dodajemy fazę stopnia mocy wyznaczoną przy wybranej częstotliwości przejścia f_c . Po zsumowaniu tych wartości wynik powinien być oddalony od granicy -360° o wielkość zapasu fazy. Możemy zatem zapisać:

$$\angle H(f_c) - 270^\circ + \text{boost} = -360^\circ + \varphi_m$$

Rozwiązując względem wartości podbicia, otrzymujemy:

$$\text{boost} = \varphi_m - \angle H(f_c) - 90^\circ$$

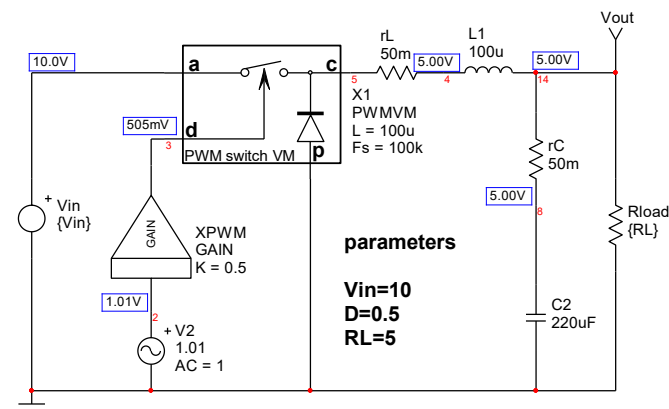
Na podstawie tej liczby wywnioskujemy, jakiego typu kompensatora użyć:

- Brak podbicia: typ 1. Odpowiedni dla przetwornic pracujących w trybie nieciągłym (DCM) oraz, do pewnego stopnia, dla stopni PFC.
- Do 90°: typ 2. Popularny w przetwornicach sterowanych prądowo (np. flyback i stopnie PFC).
- Powyżej 90° i do 180°: typ 3. Zwykle stosowany w przetwornicach sterowanych napięciowo, pracujących w trybie ciągłym przewodzenia (CCM).

Wyznaczanie odpowiedzi dynamicznej stopnia mocy

Jak już wspomniano, punktem wyjścia do analizy kompensacji danej przetwornicy impulsowej jest wykres Bodego stopnia mocy. Istnieje kilka sposobów jego uzyskania, a pierwszy z nich wykorzystuje model uśredniony w symulacji SPICE.

Modele uśrednione dostępne są w wielu odmianach, lecz najpopularniejszym jest trójkońcówkowy przełącznik PWM (three-terminal PWM switch) wprowadzony przez Vatché Vorpériana w 1986 r. i opublikowany w 1990 r. [5]. Pierwotna praca obejmowała sterowanie



Rysunek 7. Przetłacznik PWM bardzo dobrze nadaje się do symulacji uśrednionej przetwornic impulsowych, takich jak buck w tym przykładzie

napięciowe, a tryb prądowy opracowano później, lecz tylko dla CCM. W pozycji [1] wyprowadziłem samoprzełączające się (auto-toggling) wersje tych modeli zarówno dla pracy VM, jak i CM.

Typowa przetwornica buck pracująca w trybie prądowym byłaby modelowana tak, jak sugeruje **rysunek 7**. Przetłącznik PWM połączono w tzw. konfiguracji wspólnej bierno-pasywnej (common-passive), w której końcówka p jest uziemiona. Blok XPWM modeluje modulator szerokości impulsu, który przekształca napięcie błędu zadane przez źródło V_2 na współczynnik wypełnienia. Wzmocnienie tego naturalnie próbkowanego bloku modulacji jest po prostu odwrotnością wartości szczytowej piłokształtnego napięcia V_n polaryzującego komparator:

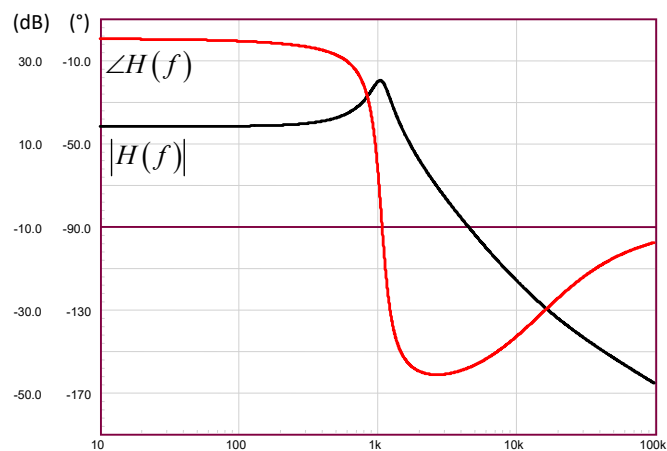
$$G_{PWM} = \frac{1}{V_p}$$

Przyjmując amplitudę szczytową piły równą 2 V, tłumienie wynosi 0,5, co odpowiada wzmacnieniu -6 dB.

Po uruchomieniu symulacji można wyświetlić punkty pracy i sprawdzić, czy są poprawne. To ważny krok pozwalający upewnić się, że przetwornica działa prawidłowo i że uzyskanym wynikom można zaufać. Tutaj model dostarcza 5 V na obciążeniu $5\ \Omega$ – i to właśnie chcieliśmy uzyskać. Wyniki można wykreślić tak, jak pokazano na **rysunku 8**.

Wypiętrzenie (peaking) w charakterystyce modułu wskazuje na duży współczynnik dobroci Q . Wielkość ta jest reprezentatywna dla strat w obwodzie i zależy od ogólnej sprawności. Jeśli zbudujesz przetwornicę buck i wykreślisz jej odpowiedź, będzie ona zapewne bardziej tłumiona niż ta z rysunku 8. Wynika to z tego, że rezystancja $R_{DS(on)}$ tranzystora MOSFET, różne straty omowe na kondensatorze i cewce oraz straty związane z odzyskiem ładunku diody zwrotnej – wszystkie przyczyniają się do strat w obwodzie i wpływają na Q .

Jeśli obciążenie wzrośnie teraz do 100 Ω , model automatycznie przechodzi w tryb DCM i dostarcza nowy wykres uzyskany przy współczynniku wypełnienia

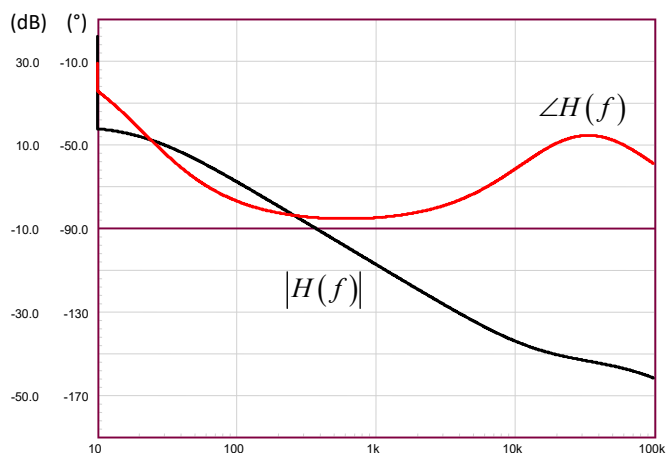


Rysunek 8. Odpowiedź drugiego rzędu wypiętrza się przy 1 kHz

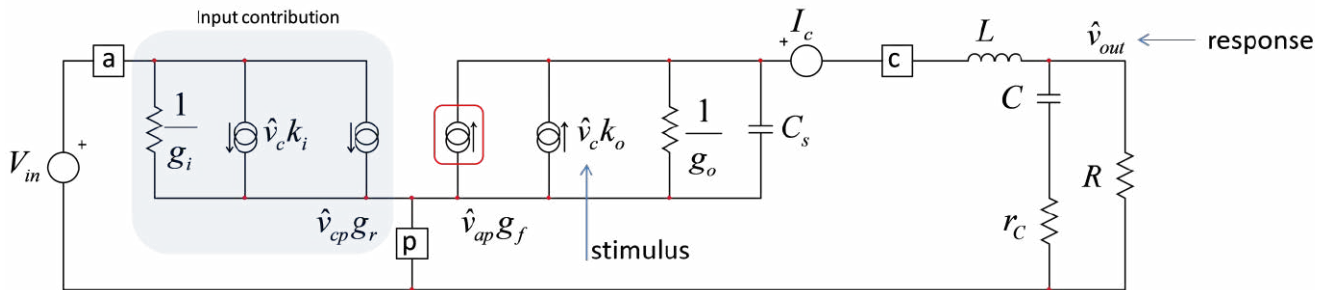
ustawionym na 31% dla tego samego napięcia wyjściowego 5 V. Zaktualizowaną odpowiedź pokazano na **rysunku 9** – potwierdza ona zanik wypiętrzenia wzmacnienia.

W odróżnieniu od wyników uzyskiwanych innymi metodami, takimi jak uśrednianie w przestrzeni stanów (SSA), przetwornica buck pracująca w DCM pozostaje układem drugiego rzędu, lecz cechuje się małym współczynnikiem dobroci Q . Widać to wyraźnie po fazie, która w modelu pierwszego rzędu zmierzałaby przy dużych częstotliwościach do zera, tymczasem dalej maleje, aż osiąga -180° . Odpowiedź składa się więc z biegunów małej i dużej częstotliwości, a kondensator wyjściowy wraz ze swoją zastępczą rezystancją szeregową (ESR) wnosi do transmisji zero.

Symulacje SPICE stanowią kolejną realną metodę wy-
 kreślenia transmitancji sterowanie–wyjście przetwornicy,
 którą chcemy ustabilizować. Choć wiernie modelują one
 wpływ pasożytów (np. ESR cewki i kondensatora), nie mó-
 wią jednak, na który człon transmitancji te pasożytni-
 cze elementy oddziałują. Zrozumienie wpływu danego ele-
 mentu na odpowiedź dynamiczną jest niezwykle istotne,
 ponieważ musisz zneutralizować jego szkodliwe skutki
 za pomocą odpowiedniej strategii kompensacji.



Rysunek 9. Pracując w trybie DCM, przetwornica buck w trybie VM nadal pozostaje układem drugiego rzędu



Rysunek 10. Przetwornica buck w trybie CM w ujęciu małosygnałowym jest opisana modelem trzeciego rzędu

Poza analizą Monte Carlo czy analizą wrażliwości, które potrafią pochłonąć wiele czasu obliczeniowego, najlepszym sposobem zrozumienia wpływu elementu na odpowiedź jest wyznaczenie transmitancji sterowanie–wyjście z modelu małosygnałowego. Taki model przedstawiono na **rysunku 10**. Tym razem wybraliśmy przetwornicę buck pracującą w trybie prądowym (CM). Jej analizę można przeprowadzić za pomocą przełącznika PWM w trybie CM, który bardzo dobrze nadaje się do tego rodzaju badań [6, 7]. Model przewiduje oscylacje subharmoniczne wynikające z niestabilnego wzmocnienia pętli prądowej. Dodając kompensację nachylenia (slope compensation), można zmniejszyć wzmocnienie pętli prądowej i skutecznie ustabilizować przetwornicę.

Zliczenie elementów magazynujących energię o niezależnej zmiennej stanu daje nam rząd tej przetwornicy: jest to układ trzeciego rzędu, a poszukujemy transmitancji sterowanie–wyjście, w której pobudzeniem jest V_c , a odpowiedzią V_{out} .

Istnieje wiele sposobów wyznaczenia zależności wiążącej V_c z V_{out} , ale moim zdaniem żaden nie dorówna szybkim technikom analitycznym obwodów (FACTs, fast analytical circuits techniques) [4]. Są one nie tylko najszybszą drogą w porównaniu z klasyczną analizą węzłową/oczkową, lecz dają również tzw. wynik o niskiej entropii (low-entropy). Licznik i mianownik w naturalny sposób przyjmują sformalizowaną postać po zakończeniu analizy. Uzyskane wyniki dają więc natychmiastowy wgląd w transmitancję: gdzie leżą bieguny i zera oraz jakie parametry na nie wpływają.

Ponownie – znajomość parametru decydującego o położeniu zera lub biegunu pozwoli skutecznie zwalczać jego naturalną zmienność w produkcji. Ray Ridley wyprowadził w swojej rozprawie doktorskiej transmitancję sterowanie–wyjście przetwornicy buck w trybie CM, uwzględniając bieguny subharmoniczne położone przy $F_{sw/2}$ [8]. Podano ją poniżej:

$$H(s) \approx H_0 \frac{1 + \frac{s}{\omega_z}}{1 + \frac{s}{\omega_p} \frac{1}{1 + \frac{s}{\omega_n Q} + \left(\frac{s}{\omega_n}\right)^2}}$$

gdzie

$$H_0 = \frac{R}{R_i} \frac{1}{1 + \frac{RT_{sw}}{L_2} [m_c(1-D)0,5]}$$

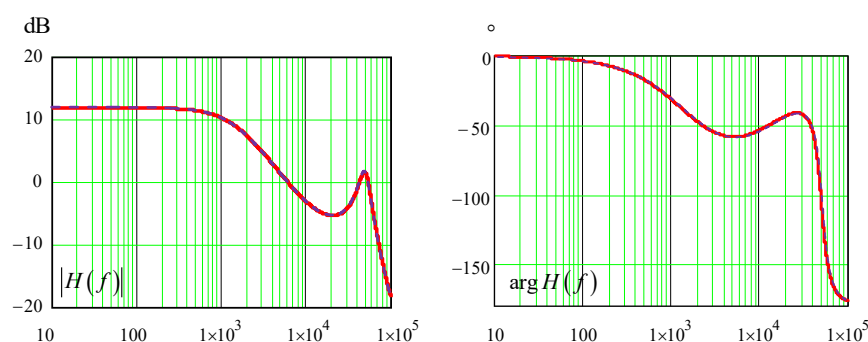
$$\omega_p = \frac{1}{RC_3} + \frac{T_{sw}}{L_2 C_3} [m_c(1-D)0,5]$$

$$\omega_z = \frac{1}{r_C C_3}$$

$$\omega_n = \frac{\pi}{T_{sw}} \quad Q = \frac{1}{\pi [m_c(1-D)0,5]}$$

W wyrażeniach tych człon m_c odnosi się do ilości zewnętrznego nachylenia celowo wstrzykniętego w modulatorze w celu zmniejszenia wzmocnienia pętli prądowej. Wielkość m_c definiuje się następująco:

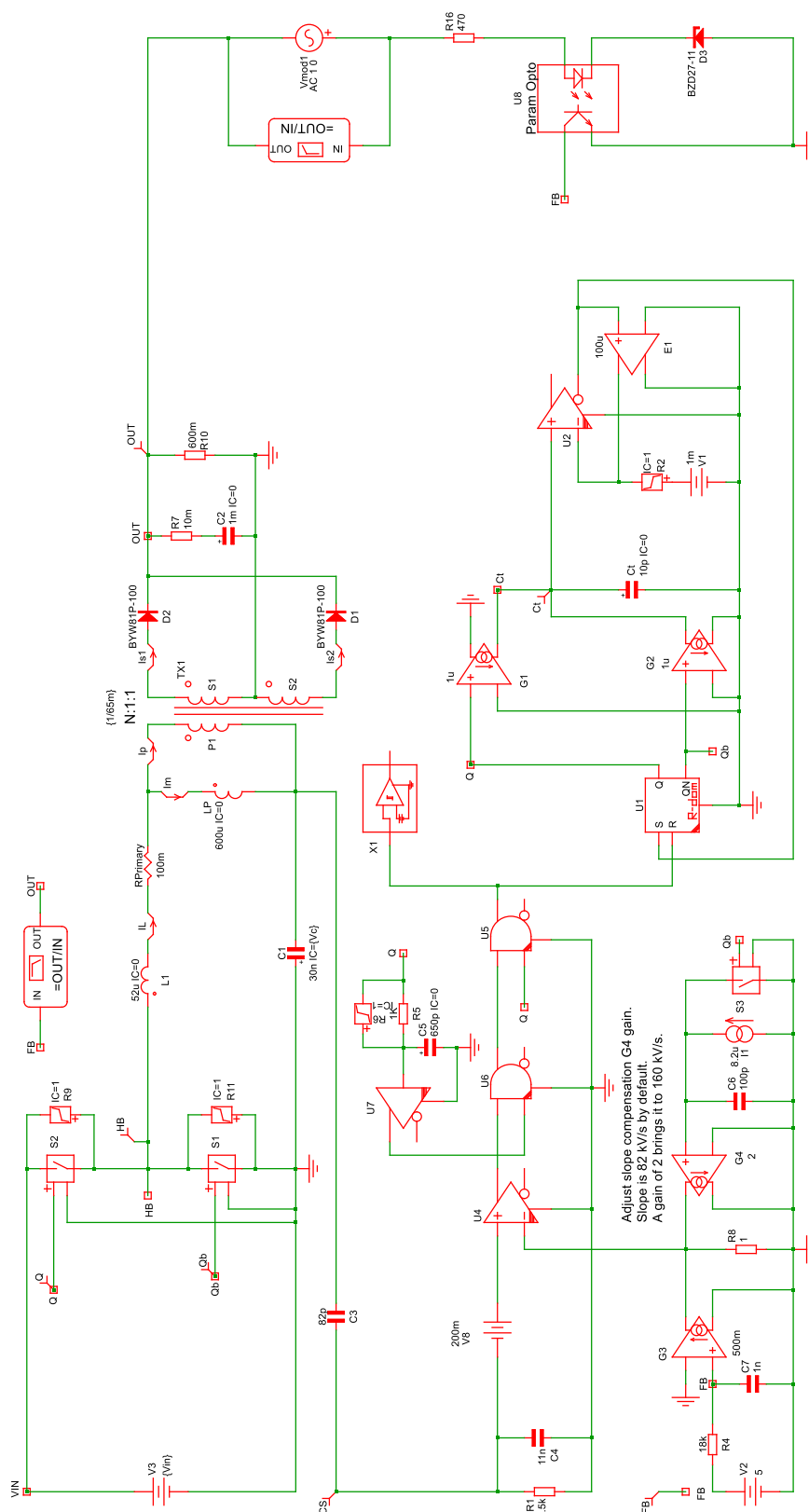
$$m_c = \frac{S_e}{S_n} + 1$$



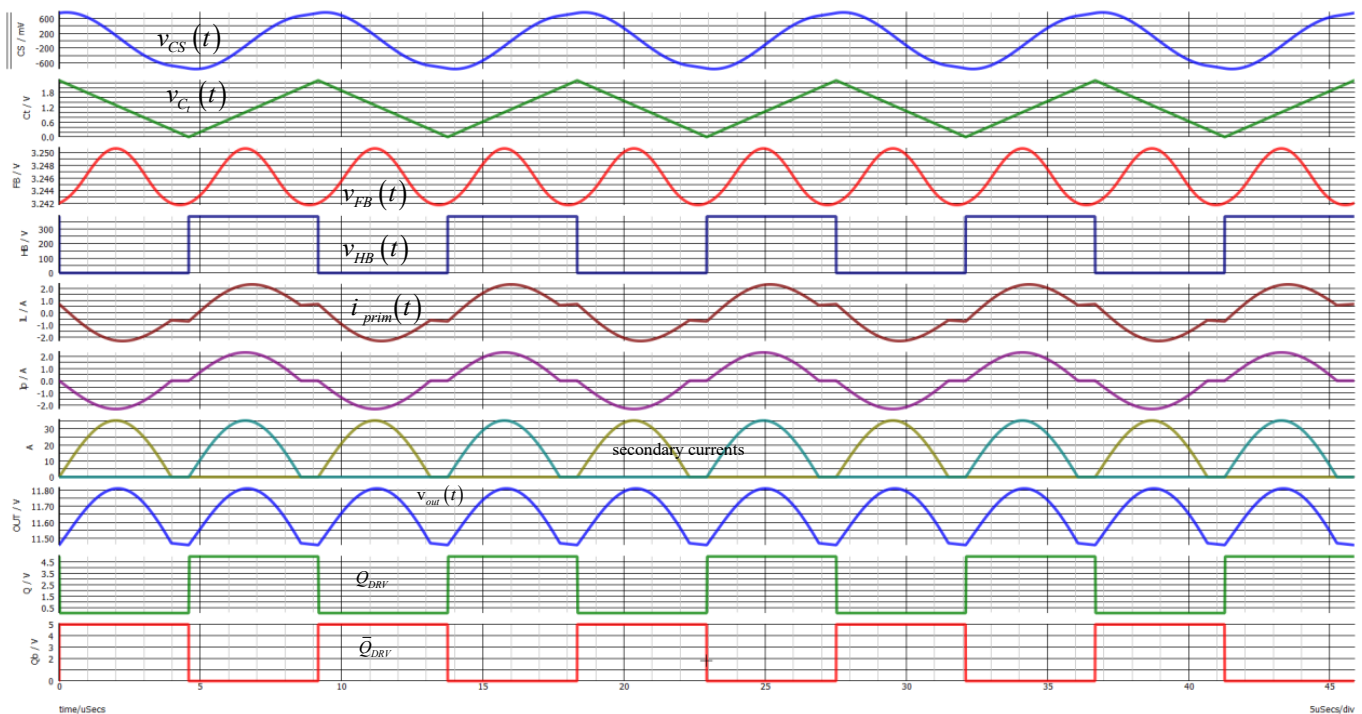
Rysunek 11. Wzmocnienie jest płaskie przy częstotliwości zerowej i opada z nachyleniem -1, aż wypiętrza się przy $F_{sw/2}$



Rysunek 12. SIMPLIS posługuje się elementami idealnymi opisanymi odcinkami liniowymi



Rysunek 13. Uproszczona wersja typowej przetwornicy LLC



Rysunek 14. Symulacja cykl po cyklu potwierdza poprawny punkt pracy – napięcie wyjściowe 24 V

S_e oznacza nachylenie zewnętrzne wyrażone w [V]/[s], natomiast S_n – nachylenie cewki w czasie załączenia, również w [V]/[s], przeskalowane przez R_i , rezystor pomiarowy. Dla przetwornicy buck nachylenie narastające prądu cewki wyznacza zależność:

$$S_n = \frac{V_{in} - V_{out}}{L_1} R_i$$

Dla $m_c = 50\%$ można wykazać, że podatność na zakłócenia wejściowe (audio susceptibility) przetwornicy buck w trybie CM jest teoretycznie zerowana.

Mając równanie (7), można wykreślić odpowiedź dynamiczną stopnia mocy i ustalić, gdzie umieścić częstotliwość przejścia. Rysunek 11 przedstawia tę odpowiedź – wyraźnie widać na niej wypiętrzenie przy połowie częstotliwości przełączania.

Widzieliśmy już, jak symulacje uśrednione oraz wyniki oparte na równaniach mogą dostarczyć potrzebnej nam odpowiedzi stopnia mocy. Trzecią opcją jest użycie symulatora zdolnego dostarczyć odpowiedź małosygnałową z obwodu przełączającego. Taki program nazywany jest symulatorem odcinkowo-liniowym (PWL, piece-wise linear).

SPICE jest w istocie solverem liniowym, a każde zachowanie nieliniowe musi zostać zlinearyzowane wokół odpowiedniego punktu pracy. Ten punkt znajdujący się przez zmniejszanie kroku symulacji aż do osiągnięcia zbieżności. Element nieliniowy, np. dioda, musi być w trakcie symulacji zastępowany przybliżeniem liniowym punktu po punkcie. Proces ten nie tylko mocno obciąża komputer, ale często rodzi błędy zbieżności, gdy algorytm redukcji kroku czasowego osiąga dolną granicę.

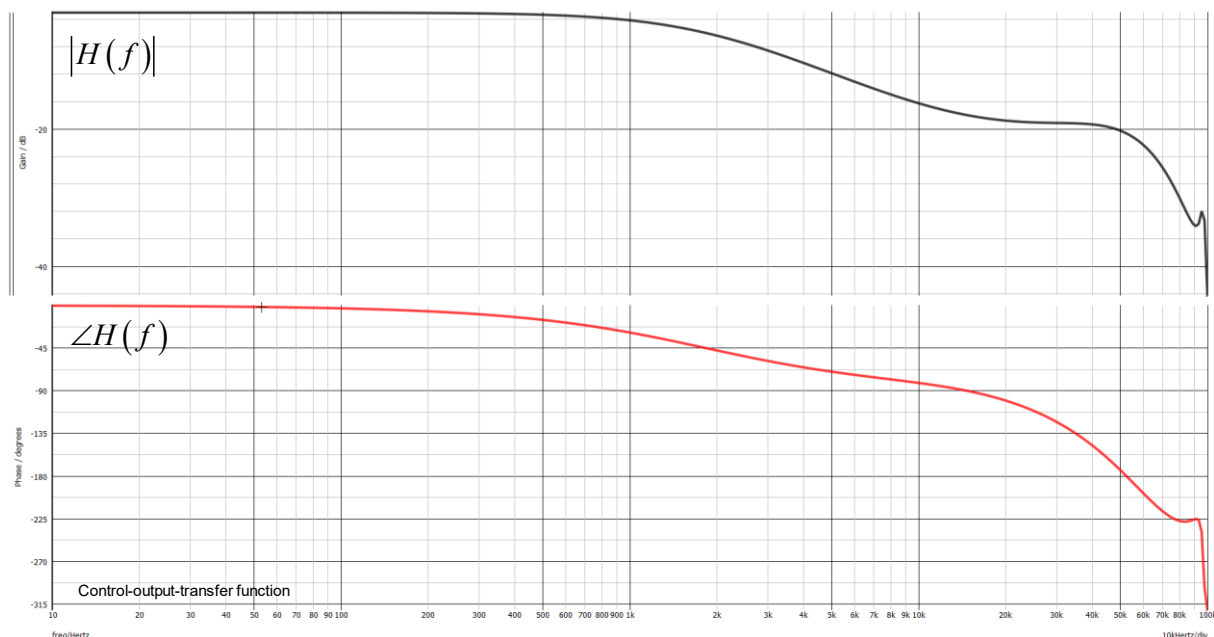
Z drugiej strony symulator taki jak SIMPLIS wykorzystuje silnik PWL i potrafi wyodrębnić odpowiedź AC z obwodu przełączającego. Rysunek 12 pokazuje, jak typowo modelowana jest dioda.

Widać, jak odcinki opisują wzrost spadku napięcia w kierunku przewodzenia w zależności od prądu diody. Z powodzeniem zastępują one wykładnicze równanie Shockleya opisujące prąd diody. Niezależnie od punktu pracy diody zachowanie jest zawsze liniowe – zmienia się jedynie nachylenie. Dzięki temu nie trzeba stosować dodatkowego algorytmu linearyzacji, ponieważ obwód jest zawsze liniowy. Można więc przyłożyć modulację AC jako pobudzenie do obwodu przełączającego i uzyskać z niego odpowiedź małosygnałową.

Typową przetwornicę LLC przedstawia rysunek 13. Tranzystory MOSFET górny i dolny pracują dokładnie przy 50% wypełnienia w tym nowym podejściu sterowania prądowego, zrealizowanym przez rezonansowy kontroler prądowy NCP13992. Tranzystor górny załącza się i pozostaje w tym stanie do chwili, gdy szczytowy prąd cewki osiągnie wartość zadaną przez pętlę sprzężenia zwrotnego.

Gdy tranzystor górny się wyłącza, aktywuje się tranzystor dolny na czas wyłączenia dokładnie odwzorowujący poprzedni czas załączenia t_{on} , co zapewnia idealne 50% wypełnienia. Zaproponowany na rysunku 13 układ jest uproszczoną wersją tej złożonej sekcji sterowania, lecz pozwala na pełną symulację obwodu za pomocą Elements – wersji demonstracyjnej programu SIMPLIS.

Po kilkudziesięciu sekundach symulator dostarcza nie tylko przebiegi cykl po cyklu – można wówczas sprawdzić np. wartości skuteczne, średnie czy szczytowe – lecz



Rysunek 15. Transmitancję sterowanie–wyjście uzyskuje się natychmiast po obliczeniu okresowego punktu pracy (POP, periodic operating point)

także transmitancję sterowanie–wyjście. Oba te wyniki przedstawiają rysunki 14 i 15.

Jest to interesujące, ponieważ nie trzeba sięgać po model uśredniony, a zarazem można badać efekty drugiego lub trzeciego rzędu, takie jak zmiany $R_{DS(on)}$, i natychmiast obserwować ich wpływ na transmitancję. Dla przetwornicy LLC istnieją modele oparte na równaniach, lecz moim zdaniem są one trudne w użyciu z uwagi na złożoność i ciężki aparat matematyczny. Dostęp do danych symulacyjnych łączących wyniki przejściowe i małosygnałowe w krótkim czasie to naprawdę użyteczne podejście.

Dobór częstotliwości przejścia i zapasu fazy

Skoro mamy już transmitancję stopnia mocy, należy wybrać i zastosować strategię kompensacji. Pierwsze pytanie brzmi: jak dobrać częstotliwość przejścia f_c oraz zapas fazy? Literatura obfituje w zalecenia mieszczące się w zakresie od jednej piątej do jednej dziesiątej częstotliwości przełączania F_{sw} . O ile górną granicą przejścia dla każdej przetwornicy jest oczywiście $F_{sw/2}$, o tyle istnieją inne ograniczenia narzucone przez przyjętą topologię.

Topologie wywodzące się z bucka

Istnieje częstotliwość rezonansowa f_0 narzucona przez obwód LC. Jeśli spojrzysz na transmitancję sterowanie–wyjście stopnia mocy w trybie napięciowym, wzmocnienie wypiętrza się przy f_0 . W konsekwencji pętla musi wykazywać pewne wzmocnienie przy tej częstotliwości, aby możliwa była korekcja ewentualnych oscylacji. Dobrą praktyką jest więc dobranie f_c na poziomie co najmniej 3- do 5-krotności częstotliwości rezonansowej.

W trybie prądowym sytuacja się upraszcza, ponieważ przy małych częstotliwościach odpowiedź ma charakter pierwszego rzędu. Może ona jednak wypiętrzać się przy $F_{sw/2}$ z powodu nietłumionych biegunów subharmonicznych. Konieczna jest wtedy kompensacja nachylenia, aby okiełznać te bieguny – przetwornica zyskuje wówczas na stabilności.

Topologie wywodzące się z buck-boosta

W tych strukturach energia przekazywana jest w procesie dwuetapowym. Najpierw magazynujesz energię w cewce w czasie załączenia, a następnie oddajesz ją do obciążenia w czasie wyłączenia. W razie nagłego zapotrzebowania na moc wyjściową przetwornica nie może zareagować natychmiast, ponieważ cewka potrzebuje kilku cykli, aby zwiększyć zmagazynowaną energię. To wrodzone opóźnienie odpowiedzi materializuje się jako zero w prawej półpłaszczyźnie (RHPZ, right-half-plane zero) w transmitancji sterowanie–wyjście.

RHPZ daje wzrost modułu (jak każde zero), lecz faza maleje. Jest to przeciwieństwo zera w lewej półpłaszczyźnie, którego faza rośnie. Gdy w transmitancji występuje RHPZ, faza stopnia mocy dodatkowo pogarsza się w miarę zbliżania się do jego położenia. Zaleca się zatem ustawienie przejścia (crossover) na dobre przed wystąpieniem RHPZ.

Dobłą zasadą jest ustawienie górnej granicy f_c na poziomie 20...30% położenia najniższego RHPZ (uzyskanego dla maksymalnego prądu i minimalnego napięcia wejściowego). Reguła ta obowiązuje zarówno dla metody VM, jak i CM, ponieważ RHPZ leży w tym samym miejscu. W trybie VM trzeba dodatkowo przestrzegać reguły bucka, czyli dobrać f_c większe niż 3- do 5-krotność f_0 – przy czym

tym razem f_0 zmienia się wraz ze współczynnikiem wypełnienia, co komplikuje ostateczny wybór.

Topologia boost

Zachowanie tej topologii jest niemal takie samo jak opisanego wyżej buck-boosta. Występuje RHPZ oraz rezonans w trybie napięciowym. W trybie prądowym jest nieco więcej swobody niż w VM, ponieważ nie ma wypiętrzenia przy f_0 , ale RHPZ i tak ogranicza górną granicę f_c . Jeśli chcesz uzyskać szerokie pasmo z przetwornicą boost lub buck-boost, lepiej zmniejszyć indukcyjność, aby przetwornica mogła szybciej reagować na nagłe zapotrzebowanie na moc wyjściową.

Wszystkie te zalecenia zebrano w tabeli 1. Pamiętaj, że zbytne podnoszenie częstotliwości przejścia tam, gdzie topologia na to pozwala, również nie jest rozsądne. Szerokie pasmo działa bowiem jak otwarcie lejka: przetwornica jest wprawdzie szybsza, lecz staje się bardziej wrażliwa na zewnętrzne zakłócenia i szumy – dobierz f_c tak, by spełnić określone wymagania dotyczące stanów przejściowych, ale nie przekraczaj tej wartości.

Dobór zapasu fazy w pętli otwartej zależy od rodzaju odpowiedzi przejściowej, jakiej potrzebujesz. Jeśli zależy Ci na szybkiej odpowiedzi i możesz zaakceptować nieco przeregulowania, zapas fazy w okolicy 50° jest odpowiedni. Jeśli wolisz zachować ostrożność i mieć wolniejszą odpowiedź (lub powrót do równowagi) bez przeregulowania, dobrym punktem wyjścia jest $70..80^\circ$. Zapas fazy w pętli otwartej φ_m można powiązać ze współczynnikiem dobroci w pętli zamkniętej Q_c za pomocą wykresu z rysunku 16. Jest to podejście teoretyczne opisujące, jak zachowuje się w pętli zamkniętej układ drugiego rzędu z biegunem w zerze i biegunem wielkiej częstotliwości (bez zer).

Trzeba zrozumieć, że dobór zapasu fazy zależy od aplikacji, ale również od granic, jakie możesz zaakceptować. Jeśli na przykład przetwornica będzie pracować w szerokim zakresie temperatur (np. od -40°C do 80°C otoczenia), lepiej wybrać duży zapas ($80..90^\circ$ lub nawet więcej)

Tabela 1. Częstotliwości przejścia nie można dobrać dowolnie – zależy ona od przyjętej topologii (założono tryb ciągłego przewodzenia, CCM)

Topologia	Tryb napięciowy (VM)	Tryb prądowy (CM)
Buck (obniżający)	$3 \cdot f_0 < f_c < F_{sw/2}^*$	$f_c < F_{sw/2}^*$
Boost (podwyższający)	$3 \cdot f_0 < f_c < 0,3 \cdot f_{RHPZ}$	$f_c < 0,3 \cdot f_{RHPZ}$
Buck-boost	$3 \cdot f_0 < f_c < 0,3 \cdot f_{RHPZ}$	$f_c < 0,3 \cdot f_{RHPZ}$

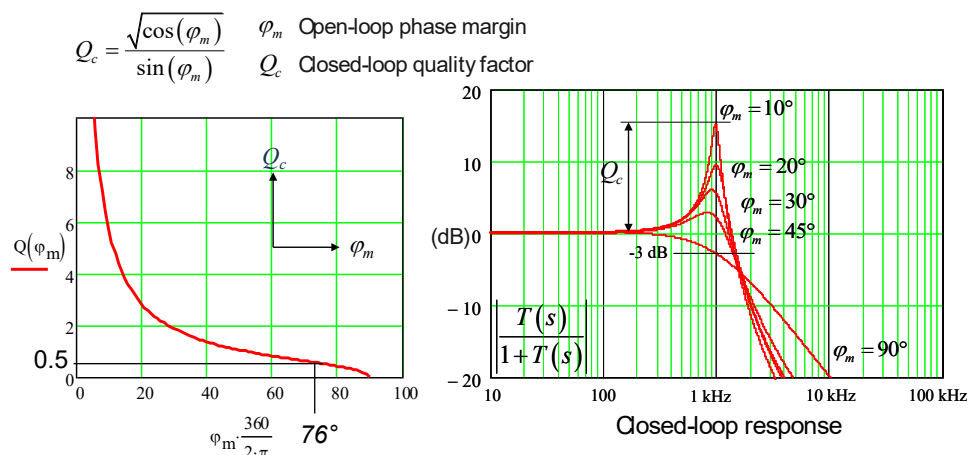
* teoretyczna granica górna

i sprawdzić, jak nisko spada on w najgorszym przypadku. Zbyt mały zapas fazy może spowodować niedopuszczalne dzwonienie odpowiedzi i ewentualne zadziałanie zabezpieczeń. Moim zdaniem 40° to rozsądna bezwzględnie najniższa wartość.

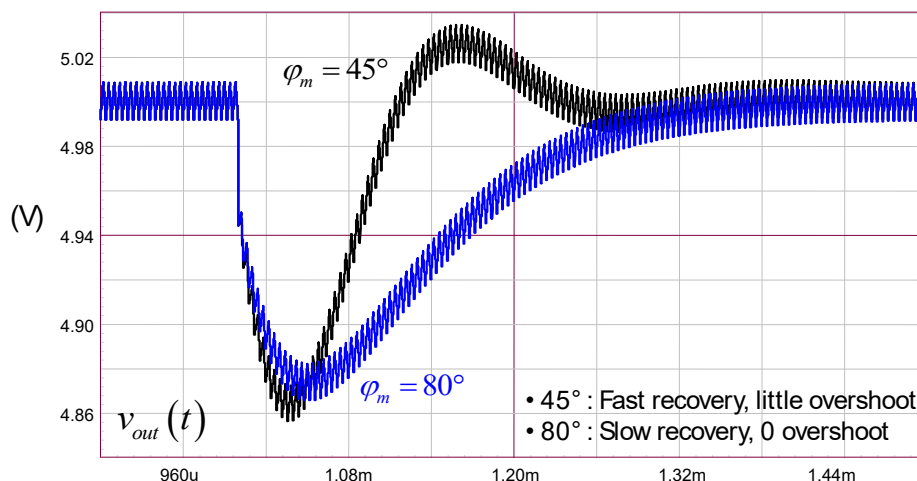
Jeśli zasilacz pracuje w pomieszczeniu, w którym temperatura otoczenia nigdy nie przekracza 35°C i nie spada poniżej 0°C (większość produktów konsumenckich), to być może łatwiej będzie spełnić mniej agresywne założenia. Po zamrożeniu projektu trzeba przeprowadzić rozległe eksperymenty (np. analizę Monte Carlo lub analizę najgorszego przypadku) i upewnić się, że w symulacjach skrajnych zapas fazy nigdy nie spada poniżej 40° .

Jak podkreśla literatura, duży zapas fazy spowolni powrót do równowagi, ale również zmniejszy wzmocnienie przy małych częstotliwościach, utrudniając przetwornicy tłumienie zakłóceń małej częstotliwości (np. tętnienia 120 Hz w przetwornicy AC-DC). Poniższy wykres pokazuje typową odpowiedź przejściową dla dwóch różnych zapasów fazy przy stałej częstotliwości przejścia (rysunek 17).

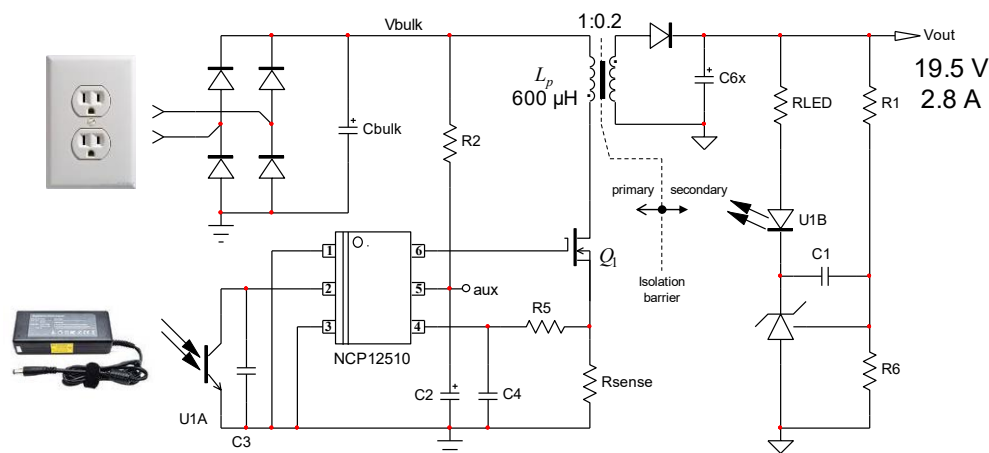
Zapas wzmocnienia zależy od zmian wzmocnienia pętli otwartej, jakich Twój układ doświadczy podczas pracy. W zależności od zmian wzmocnienia w pętli otwartej wzmacniacza błędu (spowodowanych procesami produkcyjnymi, temperaturą itd.), obecności lub braku sprzężenia w przód (feedforward) itp., moduł wzmocnienia pętli może przesunąć się w górę i w dół, wpływając na częstotliwość przejścia. Zwykle zapas wzmocnienia rzędu



Rysunek 16. Zapas fazy w pętli otwartej decyduje o tym, jak przetwornica zareaguje po zamknięciu pętli



Rysunek 17. Zbyt duży zapas fazy wpływa na czas powrotu do równowagi (f_c jest stałe)

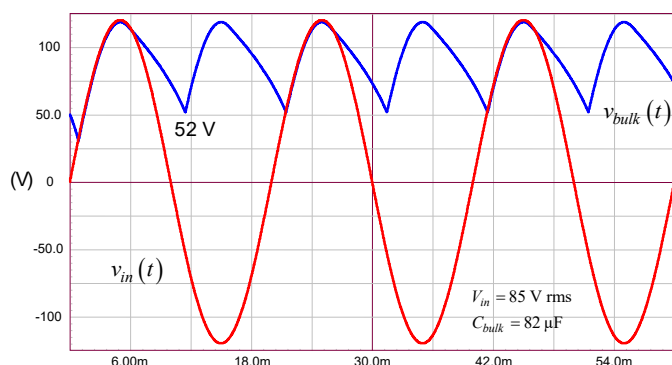


Rysunek 18. Ta przetwornica AC-DC dostarcza 2,8 A i wykorzystuje TL431 w pętli sprzężenia zwrotnego

15...20 dB uważa się za wartość zachowawczą, prowadzącą do solidnych konstrukcji [9].

Przykład projektowy: stabilizacja przetwornicy flyback AC-DC

Aby zilustrować, jak można zastosować powyższe wytyczne, wybrałem projekt zasilacza sieciowego AC-DC. Przetwornica, którą chcemy ustabilizować, dostarcza 2,8 A przy napięciu wyjściowym 19,5 V, a jej schemat przedstawia rysunek 18.

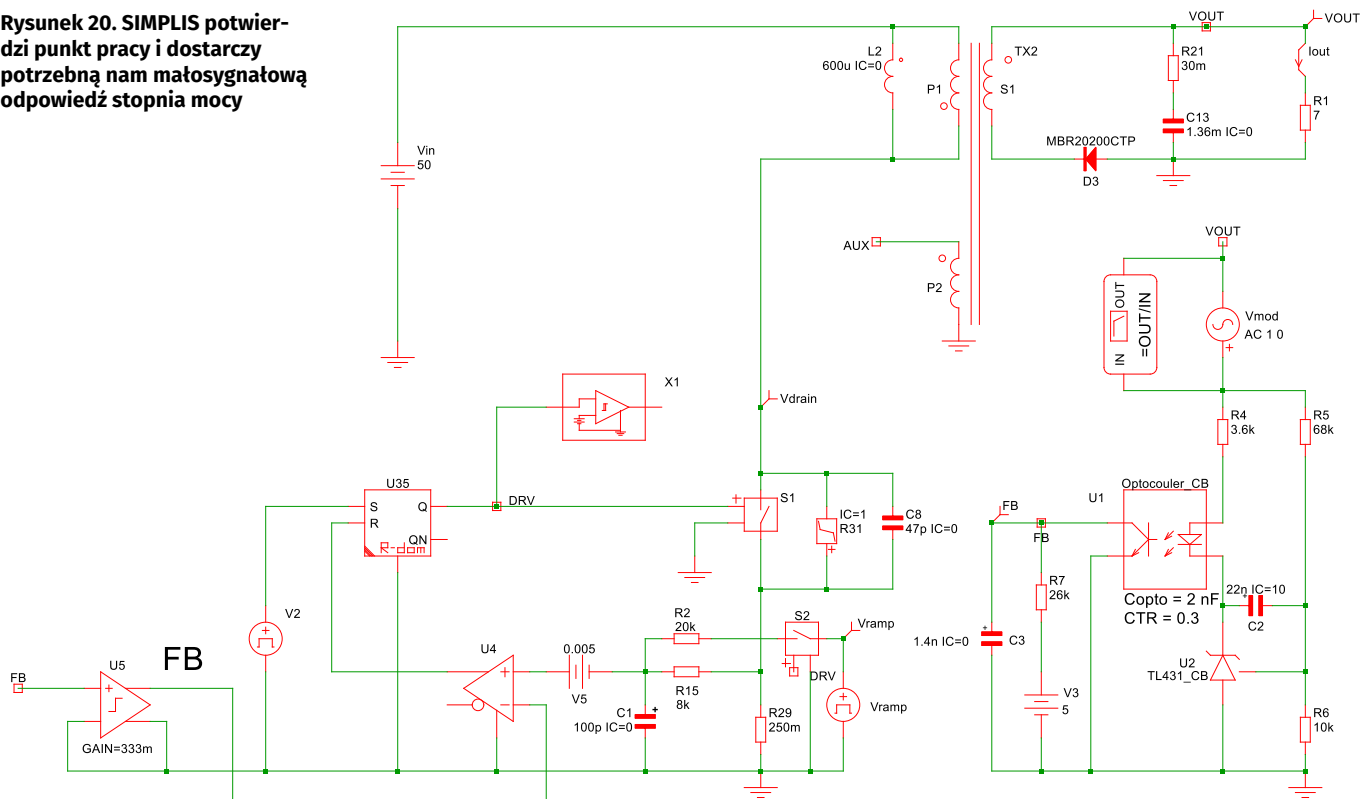


Rysunek 19. Ta przetwornica AC-DC zasilana jest z napięcia stałego obciążonego dużym tętnieniem

Ta przetwornica flyback zasilana jest z wyprostowanego napięcia sieci poprzez mostek prostowniczy i kondensator filtrujący (bulk). Z uwagi na sinusoidalne napięcie wejściowe, wyprostowane napięcie stałe zasilające flyback ma przebieg jak na rysunku 19. Widać, że przy najniższym napięciu sieci (90 V) napięcie sięga szczytowo około 120 V, lecz w dolinie spada do 52 V. Przy tym najniższym napięciu wejściowym przetwornica musi dostarczyć pełną moc, w przeciwnym razie może zadziałać zabezpieczenie przeciążeniowe lub na wyjściu pojawić się niedopuszczalne tętnienie.

Pierwszą rzeczą, jakiej potrzebujemy, jest transmitancja sterowanie–wyjście przy najniższym poziomie wejściowym. Można ją uzyskać za pomocą obwodu SIMPLIS pokazanego na rysunku 20. Bez żadnej kompensacji nachylenia spodziewamy się odpowiedzi wypiętrzającej się przy połowie częstotliwości przełączania, wynikającej z głębokiej pracy w CCM przy napięciu wejściowym 52 V. Ponieważ SIMPLIS musi odnaleźć okresowy punkt pracy (POP) przed uruchomieniem analizy AC, musiałem dodać pewną rampę kompensacyjną, aby program poprawnie się zbiegał (oprócz tej potrzebnej do zapobieżenia oscylacjom subharmonicznym).

Rysunek 20. SIMPLIS potwierdzi punkt pracy i dostarczy potrzebną nam małosygnałową odpowiedź stopnia mocy



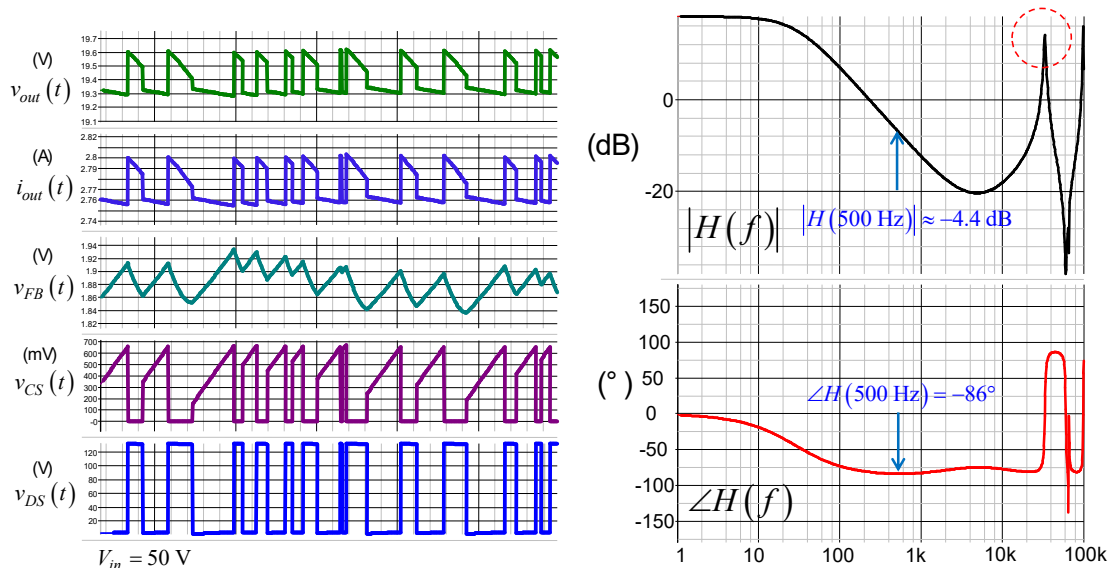
Dodatkową rampę można wygenerować z niskoimpedancyjnego pinu sterującego za pomocą obwodu RC. Jeśli stała czasowa tego obwodu jest większa od okresu przełączania, uzyskana rampa jest dość liniowa i dobrze nadaje się do celów kompensacji. Wyniki cykl po cyklu oraz wyniki AC przedstawia **rysunek 21**.

Obliczenia ujawniają ujemny współczynnik Q dla biegunów podwójnych, co wskazuje na ich położenie w prawej półpłaszczyźnie – nic dziwnego, że wzorec przełączania jest skrajnie niestabilny. Aby wybrać częstotliwość przejścia, musimy ustalić, gdzie leży zero w prawej półpłaszczyźnie. Przy tak niskim napięciu wejściowym nie zapewni nam ono dużego pasma:

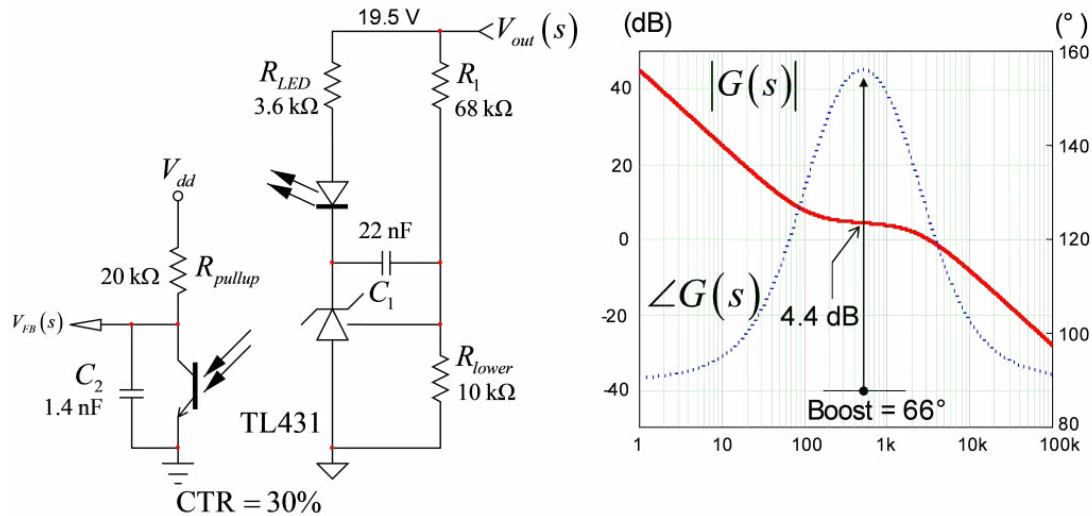
$$f_{z2} = \frac{(1-D)^2 R_{load}}{2\pi D L_p N^2} \approx 1,6 \text{ kHz}$$

gdzie N to przekładnia transformatora 1: N , D to roboczy współczynnik wypełnienia, L_p to indukcyjność uzwojenia pierwotnego transformatora, a R_{load} to rezystancja obciążenia.

Jeśli ograniczymy się do 30% tej wartości f_{z2} , to 500 Hz wydaje się rozsądną wartością f_c . Aby uzyskać lepszą wartość, masz do wyboru: albo zwiększyć kondensator filtrujący i podnieść napięcie w dolinie np. do 70 V, albo zmniejszyć indukcyjność uzwojenia pierwotnego L_p . Przesuniesz wówczas RHPZ w wyższe położenie, ale zapłacisz



Rysunek 21. Symulacja potwierdza niestabilność związaną z głęboką pracą w trybie CCM



Rysunek 22. Kompensacja typu 2 wykazuje oczekiwane podbicie fazy przy 500 Hz

za to wzrostem strat przewodzenia wynikających z większego tętnienia prądu.

Odczyt danych przy 500 Hz z wykresu Bodego stopnia mocy pokazuje tłumienie 4,4 dB i opóźnienie fazowe 86°. Możemy wyznaczyć potrzebne podbicie fazy (aby skorygować opóźnienie i wnieść pewien zapas fazy) dla celu 70° zapasu:

$$boost = \varphi_m - \angle H(f_c) - 90^\circ = 66^\circ$$

Metoda k-factor dobrze nadaje się do stabilizacji zasilaczy sterowanych prądowo i pozwala wyznaczyć, gdzie umieścić biegun i zero dla kompensatora typu 2. Najpierw wyznaczamy wartość k:

$$k = \tan\left(\frac{boost}{2} + \frac{\pi}{4}\right) = 4,7$$

Zero zostanie zatem umieszczone w:

$$f_z = \frac{f_c}{k} = \frac{500}{4,7} \approx 106 \text{ Hz}$$

natomiast biegun znajdzie się w:

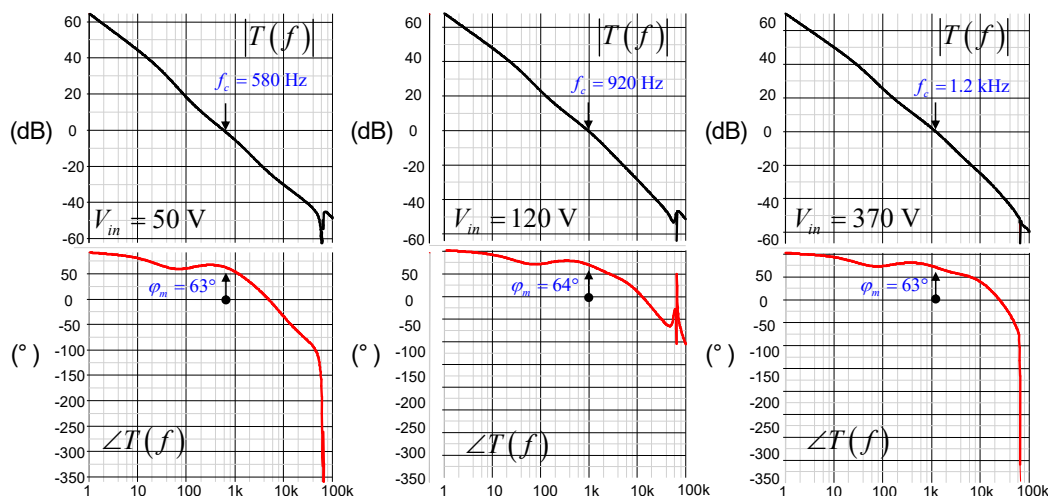
$$f_p = k \cdot f_c = 500 \times 4,7 \approx 2,35 \text{ kHz}$$

Wzmocnienie, którym chcemy skompensować niedobór 4,4 dB przy 500 Hz, zależy od rezystancji szeregowej diody LED [2] oraz współczynnika przekazu prądowego (CTR) transoptora:

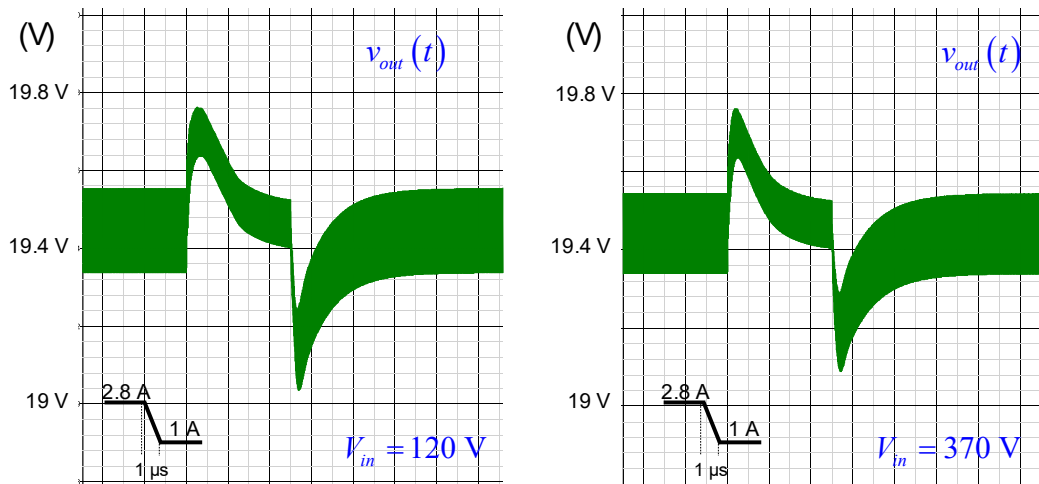
$$G_0 = \text{CTR} \frac{R_{pullup}}{R_{LED}} = 10^{\frac{-G_{fc}}{20}} = 10^{\frac{4,4}{20}} = 1,66$$

Korzystając z wartości dostępnych w naszym przykładzie projektowym, otrzymujemy rezystancję podciągającą R_{pullup} równą 3,6 kΩ. Należy sprawdzić, czy ta rezystancja jest zgodna z warunkami polaryzacji koniecznymi do sprowadzenia pinu sprzężenia zwrotnego kontrolera do stanu niskiego w najgorszym przypadku pracy.

Wreszcie biegun wielkiej częstotliwości f_p umieszcza się, dobierając pojemność równoległą do transoptora. Zwróć uwagę, że transoptor ten musi być właściwie scharakteryzowany, aby wiedzieć, gdzie leży jego biegun małej częstotliwości [2]. Po dobraniu wszystkich elementów możemy niezależnie przetestować kompensator oparty na TL431, jak pokazano na rysunku 22.



Rysunek 23. Przyjęta strategia kompensacji prowadzi do doskonałych zapasów fazy przy skrajnych napięciach wejściowych



Rysunek 24. Odpowiedź przejściowa wykazuje stabilny przebieg wyjściowy w warunkach niskiego i wysokiego napięcia sieci

Po wykonaniu tego kroku szablon SIMPLIS z rysunku 20 pozwala sprawdzić częstotliwość przejścia przy różnych napięciach wejściowych. Jak potwierdza **rysunek 23**, zapas fazy uzyskany przy skrajnych napięciach wejściowych jest bardzo komfortowy. Gdy tylko napięcie wejściowe wzrasta do 120 V, częstotliwość przejścia rozciąga się do niemal 1 kHz, co powinno korzystnie wpłynąć na szybkość reakcji.

Po wdrożeniu kompensacji przeprowadza się skoiki obciążenia na wyjściu i można sprawdzić odpowiedź. Pokazuje to **rysunek 24**. Zapad napięcia (undershoot) jest dobrze opanowany, a przy powrocie do równowagi nie ma przeregulowania. Kolejnym krokiem jest zbudowanie prototypu i weryfikacja odpowiedzi pętli na stanowisku za pomocą analizatora sieci (network analyzer).

Ten praktyczny test jest niezbędny i nie wolno go pomijać. Powie Ci on, czy założenia przyjęte podczas modelowania przetwornicy i jej układu kompensacji potwierdzają się w pomiarze na rzeczywistej płytce. Zasilanie modelu tymi danymi eksperymentalnymi pozwoli Ci przeprowadzić analizę najgorszego przypadku na komputerze i mieć pewność, że odpowiadają one rzeczywistości.

Podsumowanie

W artykule omówiono różne sposoby projektowania sekcji kompensacji przetwornicy impulsowej. Punktem wyjścia jest transmitancja sterowanie–wyjście stopnia mocy, którą można uzyskać kilkoma drogami: symulacją z modelem uśrednionym, wyprowadzeniem równań małosygnałowych lub za pomocą silnika symulacji odcinkowo-liniowej, takiego jak SIMPLIS.

Gdy działający szablon symulacyjny spełnia założone zapasy fazy i wzmocnienia, ważne jest porównanie wyników z tymi uzyskanymi na prototypie na stanowisku. Następnie na zweryfikowanym modelu przeprowadza się przemiatanie parametrów, analizę Monte Carlo oraz analizę najgorszego przypadku, aby mieć pewność, że na rynek trafia produkt solidny i niezawodny.

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

Tłumaczenie artykułu: „Analysis, Simulation And Experimentation Enable Successful Design Of Power Supply Compensation”, C. Basso, How2Power Today, lipiec 2020. ©2020 How2Power. Wszelkie prawa zastrzeżone. Etykiety i opisy na rysunkach pozostawiono w oryginalnej wersji angielskiej.

Literatura:

1. Switch-Mode Power Supplies: SPICE Simulations and Practical Designs, wyd. 2, C. Basso, McGraw-Hill, Nowy Jork, 2014.
2. „Designing Compensators for the Control of Switching Power Supplies”, C. Basso, APEC Professional Seminar, Palm Springs, 2010.
3. „The k-Factor: a New Mathematical Tool for Stability Analysis and Synthesis”, Dean Venable, Proceedings of Powercon 10, 1983.
4. Linear Circuit Transfer Functions: an Introduction to Fast Analytical Techniques, C. Basso, Wiley IEEE-Press, 2016.
5. „Simplified Analysis of the PWM Converters Using the Model of the PWM Switch, Part I (CCM) and Part II (DCM)”, Transactions on Aerospace and Electronics Systems, vol. 26, nr 3, 1990.
6. „Analysis of Current-Controlled PWM Converters Using the Model of the Current-Controlled PWM Switch”, Vatché Vorpérian, Power Conversion and Intelligent Motion Conference, 1990.
7. „Simulation and Analysis Applied to the Design of Buck Topologies”, C. Basso, APEC Professional Seminar, Anaheim, 2019.
8. „A new Small-Signal Model for Current-Mode Control”, Ray B. Ridley, rozprawa doktorska, Virginia Polytechnic Institute and State University, 1990.
9. „Loop Gain Stability Assessment”, Ray B. Ridley.

REKLAMA

facebook.com/ElektronikaPraktyczna

Przekaźniki półprzewodnikowe SSR – budowa, działanie i zastosowanie



Współczesne systemy przełączania obwodów coraz częściej rezygnują z tradycyjnych aparatów mechanicznych na rzecz technologii półprzewodnikowej. Odpowiedzią na te wymagania są przekaźniki SSR (Solid State Relay), które stanowią nowoczesną, bezstykową alternatywę dla klasycznych przekaźników elektromechanicznych (EMR).

Czym dokładnie jest przekaźnik SSR?

W ujęciu technicznym SSR to elektroniczny element przełączający, który realizuje funkcję załączania i wyłączania obwodu wykonawczego bez użycia jakichkolwiek ruchomych zestyków. Komutacja prądu odbywa się tu w sposób czysto elektroniczny. W zależności od charakteru obciążenia jako elementy wykonawcze wykorzystuje się:

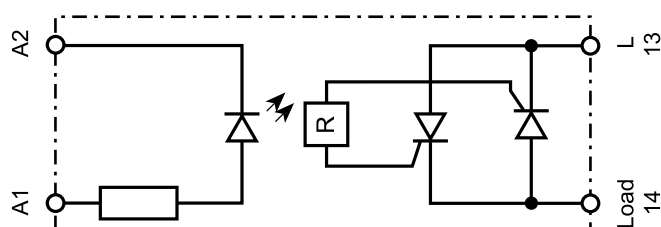
- **Triaki lub tyrystory** – w aplikacjach prądu przemiennego (AC),
- **Tranzystory (MOSFET lub IGBT)** – w układach prądu stałego (DC).

Struktura wewnętrzna typowego przekaźnika SSR opiera się na pełnej izolacji galwanicznej (najczęściej optycznej) między obwodem sterującym, a prądowym. Oznacza to, że sygnał sterujący (np. z wyjścia tranzystorowego sterownika PLC) zasila wewnętrzną diodę LED, której światło uruchamia element fotoczulawy załączający strukturę półprzewodnikową w obwodzie mocy.

Sygnał sterujący przekaźnika SSR podawany jest zazwyczaj niskim napięciem, np. 24 V DC. Wejścia sterujące dostępne są jednak w różnych zakresach napięć, takich jak 4...32 V DC, 16...32 V AC czy 90...280 V AC. W zależności od typu SSR obwód wykonawczy może natomiast przełączać napięcie przemiennego AC lub napięcie stałe DC.

Budowę SSR można przedstawić następująco:

- wejście sterujące,
- układ optoizolacji,



- układ przełączający półprzewodnikowy z wyjściem przekaźnika.

W przekaźnikach SSR nie stosuje się elementów ruchomych przez co pozbawione są one ograniczeń charakterystycznych dla klasycznych przekaźników. Przekaźnik elektromechaniczny wykorzystuje bowiem cewkę elektromagnetyczną, która po zasileniu przyciąga ruchomy styk i zamyka obwód.

Zalety SSR nad EMR

Przekaźniki SSR posiadają wiele zalet w porównaniu z klasycznymi przekaźnikami elektromagnetycznymi.

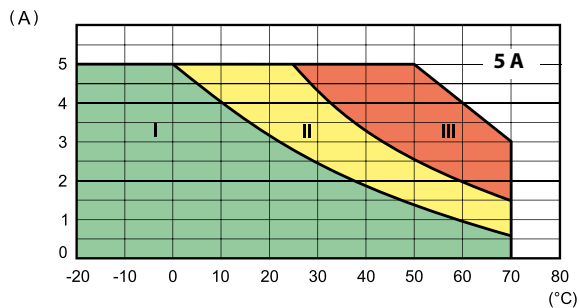
- **Bardzo duża trwałość** – Brak ruchomych elementów powoduje, że SSR mogą wykonywać miliony cykli przełączeń bez zużycia mechanicznego.
- **Cicha praca** – SSR działają całkowicie bezgłośnie, ponieważ nie występuje mechaniczne przełączanie styków.
- **Szybkie przełączanie** – Czas przełączenia jest znacznie krótszy niż w EMR, co ma znaczenie w automatyce i sterowaniu procesami.
- **Brak iskrzenia** – SSR nie powodują powstawania łuku elektrycznego, dlatego mogą pracować w środowiskach zagrożonych wybuchem lub tam, gdzie wymagane jest ograniczenie zakłóceń.

Przekaźniki elektroniczne nie są jednak pozbawione ograniczeń. Wykorzystane elementy półprzewodnikowe, podczas przewodzenia generują straty mocy i wydzielają ciepło. Z tego powodu bardzo ważny jest odpowiedni dobór radiatora oraz zapewnienie właściwego chłodzenia.

Grzanie SSR i dobór radiatora

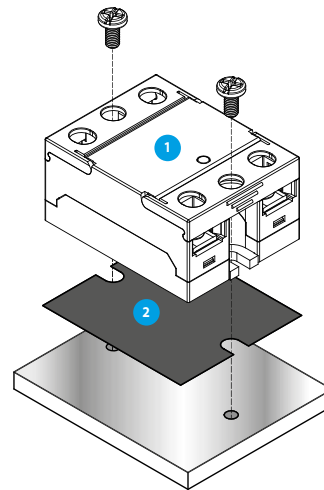
Jedną z najważniejszych cech operacyjnych przekaźnika SSR jest wydzielanie ciepła podczas pracy. Elementy półprzewodnikowe wykazują stały spadek napięcia, który

generuje straty mocy. Z tego względu przy montażu SSR należy bezwzględnie stosować radiatory oraz zadbać o odpowiednie odległości separacyjne pomiędzy modułami. Nieprawidłowe chłodzenie skutkuje przegrzewaniem urządzenia, spadkiem dopuszczalnego prądu wyjściowego, skróceniem żywotności oraz trwałym uszkodzeniem struktury półprzewodnikowej.



- I – Przełączniki zainstalowane grupowo (bez odstępu)
- II – Przełączniki zainstalowane grupowo (9 mm przerwy pomiędzy każdym)
- III – Przełączniki zainstalowane indywidualnie w wentylowanej przestrzeni (bez wpływu sąsiednich komponentów)

W kartach katalogowych przełączników SSR marki Finder znajdują się informacje na temat prawidłowej instalacji tych elementów – odległości pomiędzy SSR oraz doboru i sposobu montażu odpowiednich radiatorów.



1 Typ
77.x1
77.x2

2 Podkładka termiczna
077.T1
077.T2

NEW 077.13



77.A1
77.B1

077.T1



077.T2 - Samoprzylepne



1.3

Radiator dla SSR 1-fazowych

REKLAMA

Przełączanie o wysokiej wydajności

Seria 77

Modułowy przełącznik półprzewodnikowy (SSR)

Nowe przełączniki półprzewodnikowe typu „krążek hokejowy” do zastosowań jednofazowych aż do 125 A, dwufazowych z 2 niezależnymi kanałami aż do 75 A i do systemów trójfazowych aż do 80 A.



FINDER Polska Sp. z o.o.
ul. Logistyczna 27, 62-080 Sady
finder.pl@findernet.com

finder
RELAYING INNOVATION

findernet.com

Prądy upływu

Kolejnym istotnym ograniczeniem technologii półprzewodnikowej jest występowanie prądów upływu w stanie zablokowania przełącznika. W przeciwieństwie do klasycznych styków mechanicznych (EMR), które w stanie otwartym tworzą fizyczną przerwę powietrzną o nieskończenie wielkiej rezystancji, struktura półprzewodnikowa SSR (np. triak czy tranzystor) nawet w stanie wyłączenia przewodzą niewielki prąd rzędu kilku miliamperów (mA).

Na etapie projektowania warto dokładnie przeanalizować noty katalogowe producentów.

- SSR przeznaczone do przełączania dużych obciążeń (np. grzałek przemysłowych o prądzie 40 A) mają z reguły znacznie większe prądy upływu (nawet do 10 mA).
- Do sterowania mniejszymi elementami (np. wejściami PLC, cewkami) należy wybierać dedykowane, miniaturowe przełączniki SSR (często w obudowach interfejsowych o szerokości 6,2 mm), gdzie prąd upływu jest zredukowany do mikroamperów (μA), co całkowicie eliminuje problem.

Zastosowanie i dobór SSR

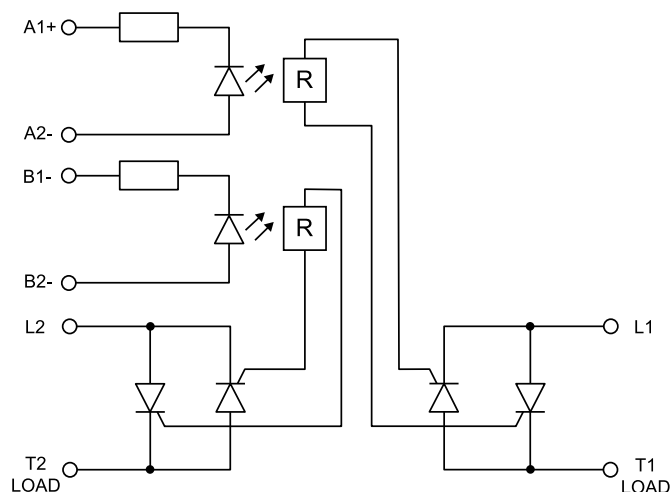
Przełączniki SSR znajdują zastosowanie przede wszystkim w automatyce przemysłowej, sterowaniu grzałkami, piecami, układami regulacji temperatury, elektrozaworami oraz innymi urządzeniami wymagającymi częstego przełączania obciążenia.

Szczególnie popularne jest wykorzystanie SSR do sterowania grzałkami elektrycznymi. W takich aplikacjach regulator temperatury może bardzo często załączać i wyłączać obciążenie. Klasyczny przełącznik EMR szybko uległby zużyciu mechanicznemu, natomiast SSR może pracować przez bardzo długi czas bez awarii.

Przy wyborze przełącznika SSR należy w pierwszej kolejności określić rodzaj i charakter obciążenia. Kategorycznie nie wolno stosować przełączników SSR AC w obwodach DC, ponieważ elementy przełączające AC mogą nie wyłączyć poprawnie prądu stałego. Dodatkowo w konfiguracjach AC dostępne są wersje z załączaniem natychmiastowym lub w chwili przejścia napięcia przez zero (Zero Crossing), co ogranicza zakłócenia sieci oraz udary prądowe przy obciążeniach o charakterze pojemnościowym.

Przełącznik S77 – wersja dwufazowa

Przełączniki typu S77 są przykładami SSR stosowanych w układach wielofazowych. Wersja dwufazowa posiada dwa niezależne zaciski przełączające, co umożliwia niezależne sterowanie dwoma fazami. Takie rozwiązanie pozwala uprościć instalację oraz zmniejszyć ilość elementów montowanych w szafie sterowniczej.



Podsumowanie

Przełączniki półprzewodnikowe SSR są nowoczesnym rozwiązaniem stosowanym szeroko w automatyce przemysłowej. Dzięki brakowi elementów mechanicznych cechują się bardzo dużą trwałością, cichą pracą oraz możliwością wykonywania ogromnej liczby przełączeń.

SSR szczególnie dobrze sprawdzają się w sterowaniu grzałkami i układami regulacji temperatury, gdzie wymagana jest wysoka niezawodność oraz częste przełączanie obciążenia.

Podczas projektowania układów z SSR bardzo ważne jest jednak odpowiednie chłodzenie, dobór radiatora oraz zachowanie właściwych odstępów montażowych. Nadmierna temperatura znacząco wpływa na skrócenie żywotności przełączników półprzewodnikowych.

Dobór odpowiedniego typu SSR – AC lub DC – powinien być zawsze dostosowany do rodzaju sterowanego obciążenia oraz parametrów pracy całego układu.

Krzysztof Smyrski



Joy-Car Calliope – robot, który uczy robotyki i programowania

Modułowy zestaw edukacyjny Joy-Car Calliope marki Joy-iT, dostępny na platformie zaopatrzeniowej Conrad, pozwala krok po kroku wprowadzać uczniów w świat robotyki – od pierwszych zajęć w szkole podstawowej po projekty realizowane na uczelniach.

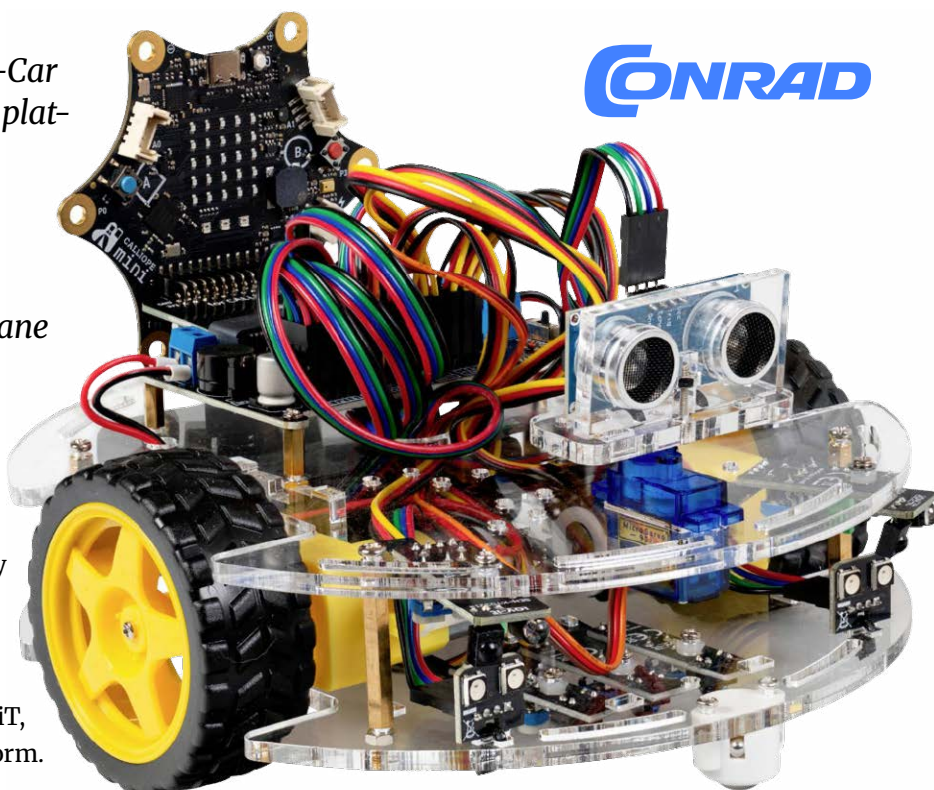
Robotyka, elektronika i programowanie należą dziś do stałych elementów programów nauczania na każdym poziomie edukacji. Skuteczne przekazywanie tej wiedzy wymaga jednak narzędzi, które łączą teorię z praktyką i potrafią rozbudzić ciekawość uczniów. Takim narzędziem jest Joy-Car Calliope – robot edukacyjny firmy Joy-iT, który trafił do oferty Conrad Sourcing Platform.

Symulacja jazdy autonomicznej

Joy-Car zbudowano na bazie platformy Calliope. Pojazd wyposażono w klakson, kierunkowskazy, reflektory oraz światła cofania i hamowania, a także rozbudowany zestaw czujników: do śledzenia linii, ultradźwiękowego pomiaru odległości, wykrywania przeszkód w podczerwieni oraz monitorowania prędkości obrotowej kół. Dzięki temu zestaw pozwala w realistyczny sposób odtworzyć zachowanie pojazdu autonomicznego, a poszczególne funkcje można programować z różnym stopniem trudności – odpowiednio do poziomu zaawansowania użytkownika.

Programowanie od dziewiątego roku życia

Robot został zintegrowany z platformą Open Roberta, opracowaną przez Instytut Fraunhofera. Wykorzystywany w niej wizualny, blokowy język NEPO zaprojektowano z myślą o nauczaniu przedmiotów ścisłych, dzięki czemu naukę mogą rozpocząć już dzieci od dziewiątego roku życia – bez wcześniejszego doświadczenia. Nauczyciele



Joy-Car to modułowy zestaw edukacyjny z bogatym wyposażeniem czujnikowym, który ułatwia start w robotyce, a zaawansowanym użytkownikom daje pole do rozbudowy
Fot. Simac Electronics GmbH/Joy-iT

i uczniowie mają do dyspozycji liczne samouczki oraz gotowe przykłady kodu. Jako zewnętrzną jednostkę sterującą zastosowano dobrze znaną w szkołach płytkę Calliope Mini.

Modułowa konstrukcja dla zaawansowanych

Joy-Car Calliope obsługuje zarówno środowiska wizualne, jak i tekstowe, dzięki czemu sprawdza się na kolejnych etapach nauki. Bardziej zaawansowani użytkownicy – na przykład w szkołach technicznych i na uczelniach – mogą sięgnąć po język MicroPython, który znacząco poszerza możliwości sterowania. Drugi, osobny moduł Calliope Mini pozwala zintegrować kierunkowskazy, oświetlenie, światła cofania i klakson oraz sterować nimi przez Bluetooth, a konfigurowalne diody LED umożliwiają tworzenie własnych efektów świetlnych.

Joy-Car Calliope marki Joy-iT jest dostępny na platformie Conrad Sourcing Platform (nr kat. 3429546). Zestaw oferowany jest również w wersji współpracującej z BBC micro:bit (nr kat. 2249432). Szczegóły i zakup: conrad.pl.

Cyfrowy nadajnik FM

Przykładowy projekt nadajnika stereo z układem KT0803L oraz przegląd cyfrowych układów nadajników FM

Nadajnik i odbiornik FM to klasyka projektów elektronicznych. Tradycyjne nadajniki na elementach dyskretnych (tranzystor, cewka, trymer) bywają jednak kapryśne: dryfują, wymagają strojenia i nie trzymają częstotliwości. Współczesna alternatywa to scalone, cyfrowo sterowane układy nadajników, które eliminują cewki i kondensatory strojeniowe. Omówienie tego tematu zawiera dwie części. W części A prezentujemy gotowy projekt – kompaktowy cyfrowy nadajnik stereo FM autorstwa Hesama Moshiriego, opublikowany w serwisie hackster.io. W części B poddajemy go krytycznej dyskusji, zwracamy uwagę na istotne kwestie prawne i zestawiamy zastosowany układ z innymi scalonymi nadajnikami FM.

Część A. Przykładowy projekt Cyfrowy nadajnik stereo FM z układem KT0803L

Nadajnik i odbiornik FM bez wątpienia należą do najpopularniejszych tematów projektów elektronicznych, jednak zbudowanie cyfrowego nadajnika FM bywa dla amatora niełatwym zadaniem. Gotowy nadajnik można podłączyć do źródła dźwięku – telefonu lub komputera – i nadawać muzykę albo inną treść audio.

W tym projekcie przedstawiono kompaktowy cyfrowy nadajnik stereo FM pracujący w paśmie częstotliwości od 87 MHz do 108 MHz. Częstotliwość przestrajają się z krokiem 0,1 MHz za pomocą dwóch przycisków chwilowych. Sercem układu jest mikrokontroler ATmega8, który komunikuje się z 0,96-calowym wyświetlaczem OLED przez interfejs SPI oraz z układem nadajnika FM KT0803L przez interfejs I²C. Do płytki można bezpośrednio podłączyć mikrofon lub kabel sygnałowy (AUX), aby nadawać wybrany dźwięk – na przykład muzykę z telefonu komórkowego czy komputera. Przeprowadzone testy wykazały, że układ pracuje stabilnie, a odbierany dźwięk jest czysty i wyraźny.

Do zaprojektowania schematu i płytki drukowanej autor użył programu Altium Designer 23, udostępniając projekt znajomym w chmurze Altium-365 w celu zebrania uwag.

Dane techniczne:

- Napięcie wejściowe: 7...9 V DC
- Pobór prądu: 50 mA
- Zakres częstotliwości: 87...108 MHz
- Krok przestrajania częstotliwości: 0,1 MHz
- Wejście dodatkowe (AUX): stereo

Autor projektu: Hesam Moshiri (współpraca: Anson Bao) – hackster.io

Tytuł oryginału: „Stereo Digital FM Transmitter Circuit (Arduino Code)”

Analiza schematu

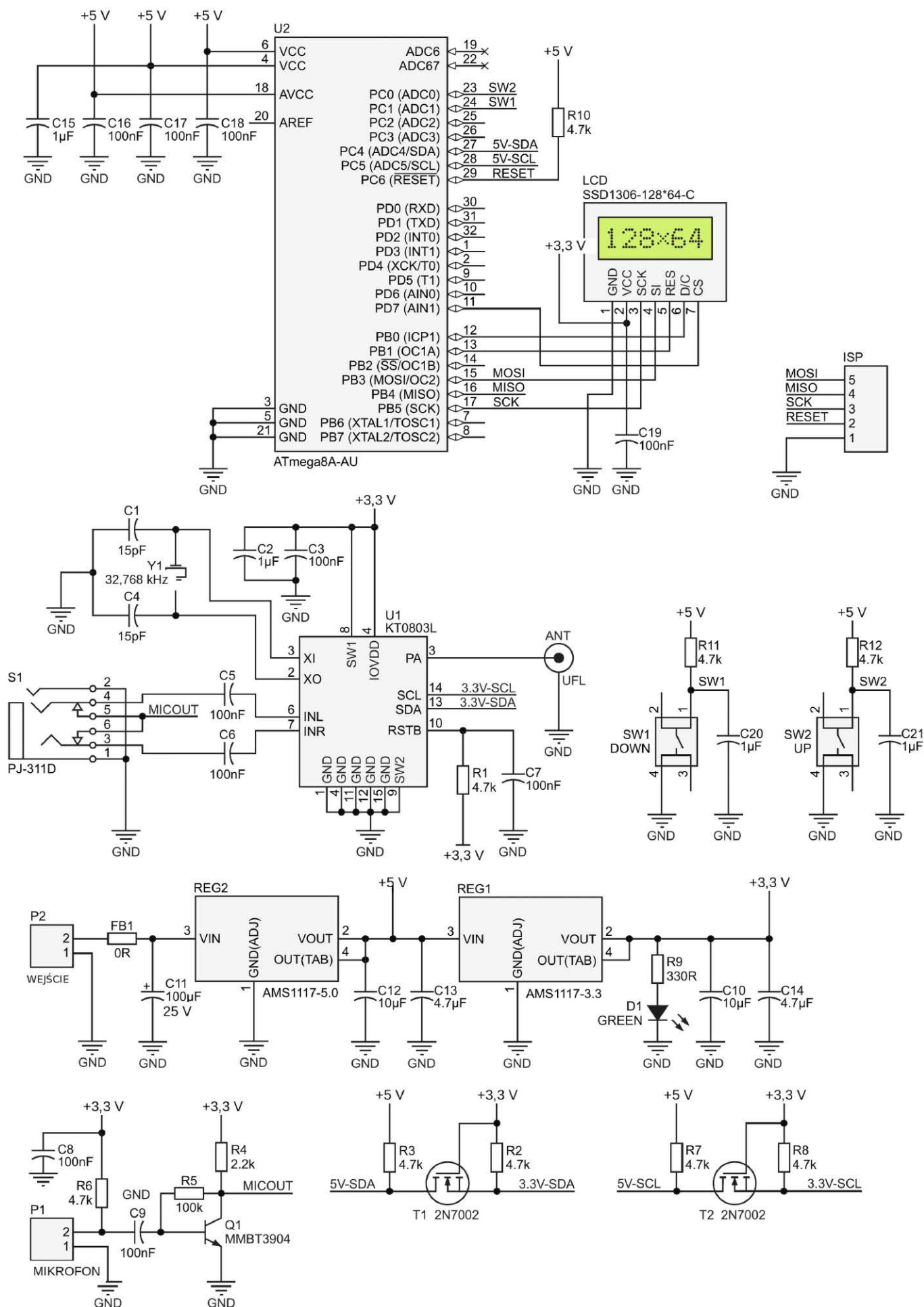
Schemat ideowy cyfrowego nadajnika FM na pasmo 87...108 MHz przedstawia **rysunek 1**. Układ składa się z kilku części, które omówiono kolejno.

Zasilanie

P2 to złącze typu XH służące do doprowadzenia zasilania do płytki; napięcie wejściowe może wynosić od 7 do 9 V DC. Elementy FB1 i C11 tworzą filtr dolnoprzepustowy ograniczający zakłócenia wnoszone przez zasilanie. REG2 to regulator TLV1117-5.0 ustalający napięcie szyny +5 V; kondensatory C12 i C13 stabilizują jego napięcie wyjściowe i ograniczają szumy. Regulator REG1 (TLV1117-3.3) ustala napięcie szyny +3,3 V. D1 to dioda LED sygnalizująca poprawne podłączenie zasilania, a kondensatory C10 i C14 stabilizują wyjście REG1 i redukują szumy.

Wejście mikrofonowe

P1 to złącze typu XH do podłączenia mikrofonu elektretowego. C8 jest kondensatorem odsprzęgającym ograniczającym zakłócenia, a rezystor R6 ustala warunki zasilania mikrofonu. Kondensator C9 usuwa składową stałą sygnału, natomiast tranzystor Q1 wzmacnia słabe sygnały mikrofonowe przed podaniem ich do układu nadajnika FM.



Rysunek 1. Schemat ideowy cyfrowego nadajnika FM (projekt w programie Altium)

Translator poziomów logicznych

T1 i T2 to tranzystory MOSFET z kanałem N typu 2N7002, służące do konwersji poziomu logicznego 5 V magistrali I²C mikrokontrolera (U2) na poziom 3,3 V wymagany przez układ nadajnika FM (U1). Rezystory R2, R3, R7 i R8 są rezystorami podciągającymi (pull-up) uzupełniającymi obwód translatora.

Nadajnik FM

Głównym elementem tej części układu jest scalony nadajnik KT0803L (U1). S1 to gniazdo słuchawkowe w wykonaniu SMD, służące do podłączenia kabla AUX – pozwala ono przesyłać dźwięk z telefonu, komputera lub innego urządzenia do płytki. Kondensatory C5 i C6 przekazują sygnał akustyczny do układu U1. C2 i C3 są kondensatorami odsprężającymi, a ANT to złącze wielkiej częstotliwości (UFL) umożliwiające podłączenie do płytki anteny teleskopowej.

Mikrokontroler i wyświetlacz

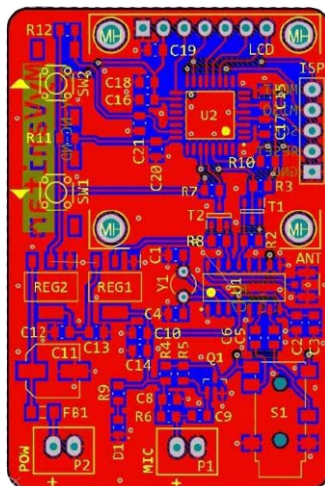
Sercem układu jest mikrokontroler ATmega8-AU (U2). C15...C18 to kondensatory odsprężające ograniczające poziom zakłóceń. R10 jest rezystorem podciągającym wejście RESET. LCD to 0,96-calowy wyświetlacz OLED o rozdzielczości 128×64 z interfejsem SPI, pokazany na **rysunku 2**. C19 to kondensator odsprężający wyprowadzenie zasilania wyświetlacza. SW1 i SW2 to przyciski chwilowe służące do zwiększania i zmniejszania częstotliwości. Kondensatory C20 i C21 eliminują drgania styków przycisków, a R11 i R12 są rezystorami podciągającymi. Złącze ISP udostępnia wyprowadzenia AVR-ISP do zaprogramowania mikrokontrolera; choć nie jest to obowiązkowe, można w jego miejscu włutować listwę kołkową.

Płytką drukowaną

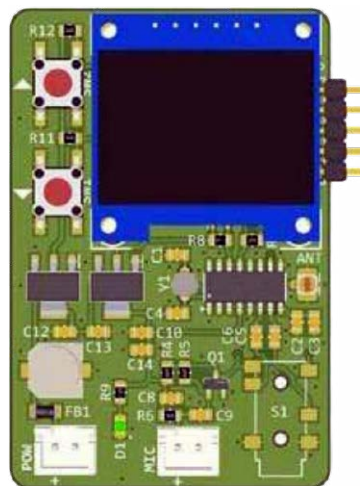
Mozaikę ścieżek płytki drukowanej przedstawiono na **rysunku 3**. Płytką jest dwustronna, a niemal



Rysunek 2. Żółto-niebieski wyświetlacz OLED SPI 0,96" o rozdzielczości 128×64



Rysunek 3. Mozaika ścieżek płytki drukowanej cyfrowego nadajnika FM



Rysunek 4. Rozmieszczenie elementów na płycie drukowanej nadajnika FM

wszystkie zastosowane elementy są w wykonaniu SMD. Rozmieszczenie elementów na płycie pokazuje **rysunek 4**.

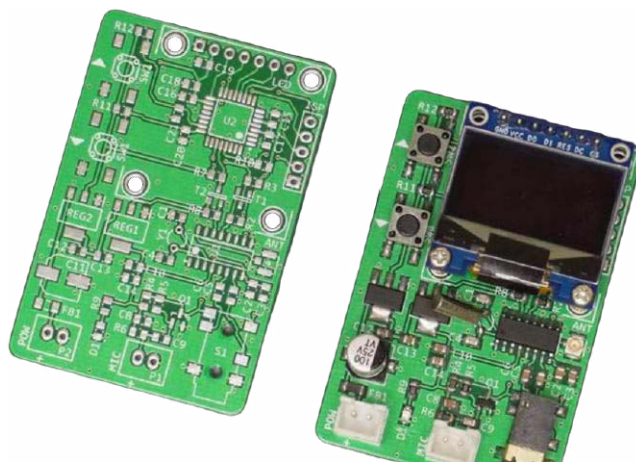
Wykaz elementów

W oryginale autor dołączył wykaz elementów wygenerowany w serwisie Octopart. Zestawiono go w formie tabeli (**tabela 1**), na podstawie oznaczeń ze schematu ideowego (**rysunek 1**). Uwaga redakcji: numer katalogowy mikrokontrolera podany w odsyłaczach oryginału (ATmega8U2) dotyczy innego układu niż użyty na schemacie (ATmega8-AU). W tabeli przyjęto oznaczenie zgodne ze schematem – ATmega8A-AU.

Program i programowanie

Jeśli układ ma zostać zbudowany dokładnie według projektu, wystarczy pobrać plik HEX (z sekcji materiałów do pobrania) i zaprogramować mikrokontroler. Bity konfiguracyjne (Fuse) należy ustawić na wewnętrzne źródło zegara 8 MHz.

Aby zmodyfikować program, autor udostępnił kod dla platformy Arduino. Do środowiska Arduino IDE należy dołączyć bibliotekę FM (obsługa KT0803L),



Rysunek 5. Zmontowana płytka cyfrowego nadajnika FM

bibliotekę wyświetlacza SPI OLED oraz menedżer płytek MiniCore. Źródło zegara ustawia się na wewnętrzne, 8 MHz. Szkielet kodu sprowadza się do zainicjowania nadajnika i wyświetlacza w funkcji `setup()`, po czym pętla `loop()` odpytuje dwa przyciski i – po wykryciu naciśnięcia – zmienia częstotliwość o 0,1 MHz, wywołuje funkcję ustawienia częstotliwości nadajnika i odświeża wskazanie na wyświetlaczu OLED. Krótkie opóźnienie po wykryciu naciśnięcia pełni rolę programowej eliminacji drgań styków.

Tabela 1. Wykaz elementów nadajnika FM

Oznaczenie	Element	Uwagi
U1	KT0803L	scalony nadajnik FM, interfejs I ² C
U2	ATmega8A-AU	mikrokontroler AVR, obudowa TQFP-32
LCD	OLED SSD1306 128×64	wyświetlacz 0,96", interfejs SPI
REG1	TLV1117-3.3 (AMS1117-3.3)	regulator LDO szyny +3,3 V
REG2	TLV1117-5.0 (AMS1117-5.0)	regulator LDO szyny +5 V
Q1	MMBT3904	tranzystor NPN – wzmacniacz mikrofonowy
T1, T2	2N7002	MOSFET N – translator poziomów I ² C
D1	dioda LED zielona	sygnalizacja zasilania
Y1	rezonator kwarcowy 32,768 kHz	źródło zegara odniesienia KT0803L
C1, C4	15 pF	kondensatory rezonatora kwarcowego
C2, C3, C7... C10, C15... C19	100 nF	kondensatory odsprężające
C5, C6	100 nF	sprzężenie sygnału audio do U1
C11	100 μF/25 V	kondensator wejściowy zasilania
C12, C13	10 μF/4,7 μF	filtrowanie wyjścia REG2
C14	4,7 μF	filtrowanie wyjścia REG1
C20, C21	1 μF	eliminacja drgań styków SW1/SW2
R1...R3, R6... R8, R10... R12	4,7 kΩ	rezystory podciągające i polaryzujące
R4	2,2 kΩ	obwód wzmacniacza mikrofonowego
R5	100 kΩ	polaryzacja bazy Q1
R9	330 Ω	rezystor szeregowy diody LED D1
FB1	koralek ferrytowy (0 Ω)	filtr zakłóceń wejścia zasilania
SW1, SW2	przyciski chwilowe (tact switch)	regulacja częstotliwości
P1, P2	złącza XH 2-pin	mikrofon/zasilanie
S1	gniazdo słuchawkowe SMD (PJ-311D)	wejście AUX
ANT	złącze UFL	antena teleskopowa
ISP	listwa kołkowa 2×3	programowanie AVR-ISP

Montaż i uruchomienie

Lutowanie elementów nie powinno sprawić trudności, ponieważ najmniejsza zastosowana obudowa to 0805, a jedynym wymagającym elementem jest złącze UFL. Osobom, które nie mają czasu lub doświadczenia w montażu, autor sugeruje zamówienie płytki w wersji zmontowanej. Zmontowaną płytkę przedstawia rysunek 5.

Część B. Dyskusja i projekty pokrewne

Część B składa się z dwóch wątków. Najpierw poddajemy przykładowy projekt krytycznej ocenie – wskazujemy istotne ograniczenia, zwłaszcza prawne, oraz miejsca warte uzupełnienia. Następnie umieszczamy zastosowany układ KT0803L w szerszym kontekście: przeglądamy rodzinę scalonych nadajników FM, od najprostszych modulatorów po układy z RDS i regulacją cyfrową, wraz z zestawieniem porównawczym.

B.1. Kwestia zasadnicza – legalność nadawania

WAŻNE

– ograniczenia prawne nadawania w paśmie FM

Pasma 87,5...108 MHz jest pasmem radiofonii UKF i podlega ścisłej regulacji. W Unii Europejskiej, w tym w Polsce, nadajniki FM małej mocy mogą pracować bez pozwolenia radiowego wyłącznie wtedy, gdy ich moc promieniowana zastępcza (ERP) nie przekracza 50 nW – zgodnie ze zharmonizowaną normą ETSI EN 301357 oraz zaleceniami CEPT (por. raport EBU R120).

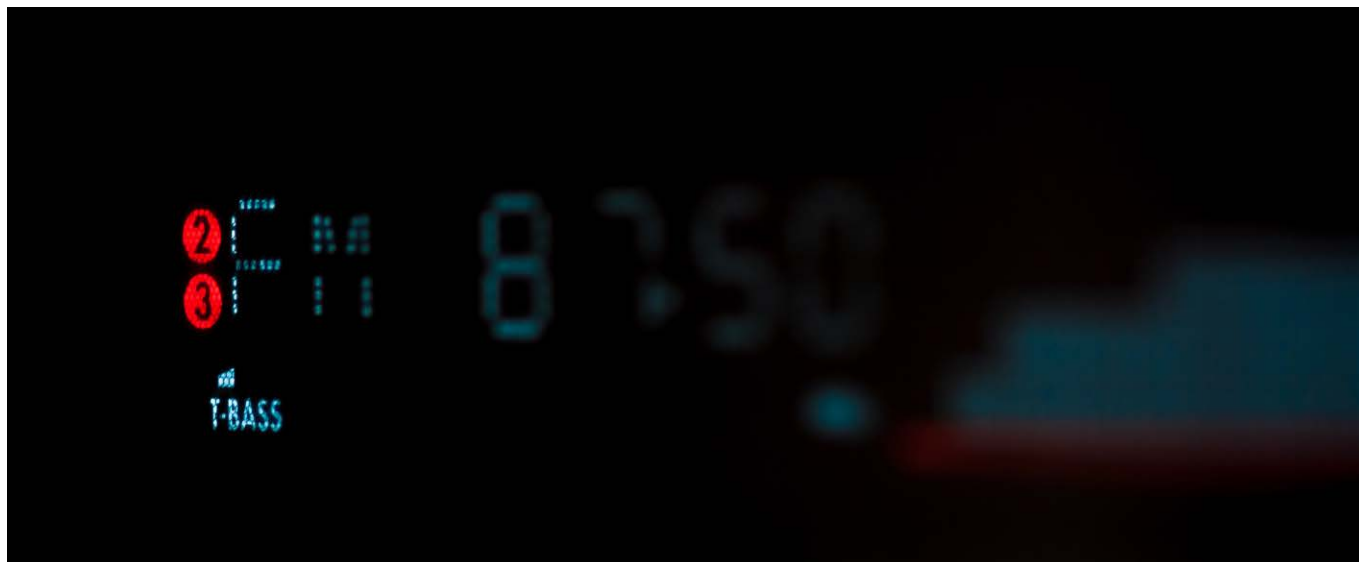
Limit 50 nW jest bardzo niski – odpowiada zasięgowi rzędu pojedynczych metrów, wystarczającemu np. do przesłania dźwięku do radioodbiornika w tym samym pomieszczeniu lub w samochodzie. Układ KT0803L może wytwarzać sygnał znacznie silniejszy, a podłączenie anteny teleskopowej (złącze ANT w projekcie) dodatkowo zwiększa moc wypromieniowaną – łatwo wówczas przekroczyć dopuszczalny limit.

Nadawanie z mocą ponad dozwolony limit, na zajętej częstotliwości lub z anteną zewnętrzną może stanowić nielegalną emisję radiową, powodować zakłócenia i podlegać karom. Projekt należy traktować jako układ edukacyjny do pracy z minimalną mocą i anteną zastępczą.

B.2. Komentarze do przykładowego projektu

Projekt Hesama Moshiriego to staranna, kompaktowa konstrukcja oparta na sprawdzonym, scalonym nadajniku FM. Jest atrakcyjny prostotą – układ KT0803L eliminuje cewki i elementy strojeniowe, a cyfrowe sterowanie zapewnia stabilną częstotliwość. Poniższe uwagi wskazują miejsca, w których oryginalny opis warto doprecyzować lub w których konstrukcję dałoby się ulepszyć.

- Brak filtru wyjściowego i dopasowania anteny.** Schemat prowadzi sygnał wyjściowy z układu KT0803L wprost do złącza UFL i anteny. Scalone nadajniki FM tej klasy wytwarzają jednak zauważalny



poziom harmonicznych powyżej pasma użytkowego. W konstrukcji nadającej się do realnego użytku warto przewidzieć prosty filtr dolnoprzepustowy LC na wyjściu antenowym – ogranicza on emisje niepożądane i poprawia widmową czystość sygnału. Brakuje też obwodu dopasowania impedancji do konkretnej anteny; w wersji edukacyjnej pracującej z minimalną mocą nie jest to krytyczne, ale przy jakiegokolwiek rozbudowie staje się istotne.

2. Dobór mikrokontrolera – ATmega8 to konstrukcja archaiczna. ATmega8 jest układem sprawdzonym, lecz przestarzałym i stosunkowo drogim w przeliczeniu na możliwości. Całe zadanie – obsługa dwóch przycisków, magistrali I²C oraz wyświetlacza SPI – bez trudu realizuje dowolny nowszy mikrokontroler AVR (np. z rodziny ATmega48/88/168) lub tani układ z rdzeniem ARM Cortex-M0. Nowsze AVR mają ponadto sprzętowy interfejs I²C lepiej udokumentowany w bieżących narzędziach. Dla powtarzalności projektu nie jest to wada, ale przy własnych modyfikacjach warto rozważyć nowszą podstawę.

3. Niespójność oznaczeń w dokumentacji źródłowej. W odsyłaczach oryginału mikrokontroler opisano numerem ATmega8U2 (układ z interfejsem USB), podczas gdy schemat i opis jednoznacznie wskazują ATmega8-AU (klasyczny AVR w obudowie TQFP-32). To rozbieżność w dokumentacji, nie w samym układzie – przy zamawianiu elementów należy kierować się oznaczeniem ze schematu. W tabeli elementów (tabela 1) przyjęto poprawne oznaczenie ATmega8A-AU.

4. Translacja poziomów I²C – rozwiązanie poprawne, lecz wymaga uwagi przy rezystorach. Zastosowanie dwóch tranzystorów 2N7002 jako translatora poziomów 5 V/3,3 V dla magistrali I²C to klasyczne, poprawne rozwiązanie (układ Philips/NXP AN10441). Jego prawidłowe działanie zależy

jednak od rezystorów podciągających po obu stronach – przyjęta wartość 4,7 k Ω jest rozsądna dla I²C w trybie standardowym, ale przy długich ścieżkach lub większej pojemności magistrali może wymagać korekty. Warto też pamiętać, że dren tranzystora musi być podłączony do strony 5 V, a źródło do strony 3,3 V – odwrotne podłączenie uniemożliwia działanie.

5. Zasilanie 7...9 V przy poborze 50 mA – straty ciepłe na regulatorach. Płytkę zasilają napięciem 7...9 V, a obie szyny (5 V i 3,3 V) tworzą regulatory liniowe LDO. Przy napięciu wejściowym 9 V i poborze 50 mA na regulatorze 5 V odkłada się ok. 4 V, co daje ok. 0,2 W strat – niewiele, lecz w obudowie SOT-223 bez radiatora warto to uwzględnić. Zasilanie układu napięciem bliższym dolnej granicy (7 V) zmniejsza straty i nagrzewanie. Alternatywą – przy własnych modyfikacjach – jest zasilanie pojedynczą szyną 5 V z portu USB i pojedynczy regulator 3,3 V.

6. Rezonator 32,768 kHz jako zegar odniesienia. KT0803L korzysta z zegara odniesienia; w projekcie zastosowano rezonator zegarkowy 32,768 kHz z kondensatorami obciążającymi 15 pF (C1, C4). To rozwiązanie zgodne z notą katalogową układu. Należy zadbać o krótkie ścieżki rezonatora i poprawne masy – przy nieostrożnym rozmieszczeniu zegar odniesienia bywa źródłem niestabilności częstotliwości nadawania.

7. Wejście mikrofonowe i wejście AUX dzielą tę samą drogę sygnału. Mikrofon elektretowy (wzmacniany tranzystorem Q1) oraz wejście AUX trafiają ostatecznie do tych samych wejść audio układu U1. W praktyce oznacza to, że jednoczesne podłączenie obu źródeł spowoduje ich zmieszanie. Nie jest to wada – projekt zakłada korzystanie z jednego źródła naraz – ale warto o tym wiedzieć; rozbudowa o przełącznik źródeł lub prosty sumator audio bywa pożądana.

8. Brak regulacji dewiacji i poziomu wejściowego. KT0803L udostępnia (przez I²C) ustawienia poziomu wejściowego i siły nadawania, lecz program przykładowy ogranicza się do przestrajania częstotliwości. Przy zbyt silnym sygnale audio modulacja może być przesterowana, co słychać jako zniekształcenia. Rozbudowa oprogramowania o regulację wzmocnienia wejściowego (a także o funkcje takie jak wyłączanie nadajnika czy odczyt stanu) wykorzystaby możliwości układu, których projekt w obecnej wersji nie używa.

Podsumowanie. Projekt jest dobrym, zwartym przykładem zastosowania scalonego nadajnika FM i świetnie nadaje się do nauki obsługi magistral I²C i SPI. Najważniejsze zastrzeżenie ma charakter prawny (punkt B.1) – przy jakimkolwiek praktycznym użyciu trzeba bezwzględnie respektować limit mocy. Wskazane uzupełnienia (filtr wyjściowy, regulacja dewiacji, ewentualnie nowszy mikrokontroler) podnoszą jakość i funkcjonalność konstrukcji.

B.3. Scalone nadajniki FM – przegląd

Zastosowany w projekcie KT0803L należy do szerokiej rodziny scalonych nadajników FM. Wszystkie te układy zastępują tradycyjny generator LC i tor wielkiej częstotliwości pojedynczą strukturą krzemową, różnią się jednak stopniem integracji, sposobem sterowania i dostępnością dodatkowych funkcji, takich jak kodowanie stereo czy transmisja danych RDS.

KT0803L (KT Micro)

Monolityczny cyfrowy nadajnik stereo FM sterowany przez I²C. Ma wbudowany przetwornik audio, cyfrowy procesor sygnału stereo oraz tor wielkiej częstotliwości; wymaga jedynie zegara odniesienia i minimum elementów zewnętrznych. To właśnie ten układ jest podstawą przykładowego projektu – jego zaletą jest prostota, mała obudowa i pełna kontrola cyfrowa.

QN8027 (Quintic)

Nowszy cyfrowy nadajnik FM sterowany przez I²C, ceniony w środowisku konstruktorów za stabilność częstotliwości i dobrą jakość dźwięku. Obsługuje kodowanie stereo i bywa wskazywany jako rozsądny wybór wszędzie

tam, gdzie zależy nam na powtarzalności i krótkim czasie uruchomienia.

Si4713 (Skyworks, dawniej Silicon Labs)

Zaawansowany cyfrowy nadajnik FM sterowany przez I²C, z obsługą transmisji danych RDS/RBDS oraz pomiarem zajętości kanału. Jest najbardziej rozbudowanym układem z tego zestawienia – kosztem wyższej ceny i bardziej złożonego oprogramowania. To naturalny kierunek rozbudowy projektu o wyświetlanie nazwy stacji w odbiorniku.

BH1417 (ROHM)

Stereofoniczny nadajnik FM z pętlą PLL, sterowany sprzętowo – wybór jednej z 14 częstotliwości następuje przełącznikiem DIP, a nie magistralą cyfrową. Zapewnia separację kanałów ok. 40 dB. Wymaga rezonatora kwarcowego (nominalnie 7,6 MHz; w praktyce stosuje się łatwiej dostępny 7,68 MHz). To rozwiązanie pośrednie między układem analogowym a w pełni cyfrowym.

BA1404 (ROHM)

Klasyczny, analogowy monolityczny nadajnik stereo FM z wbudowanym modulatorem stereo, modulatorem FM i wzmacniaczem wielkiej częstotliwości. Zapewnia nieco lepszą separację kanałów (ok. 45 dB) niż BH1417, lecz jako układ analogowy nie oferuje cyfrowego ustawiania częstotliwości – strojenie odbywa się obwodem LC. Bywa traktowany jako układ „klasyczny”, dziś raczej wypierany przez rozwiązania cyfrowe.

B.4. Zestawienie porównawcze układów

Tabela 2 zestawia omówione układy scalone według sposobu sterowania i kluczowych cech. Przykładowy projekt z części A opiera się na układzie KT0803L – w pełni cyfrowym, sterowanym przez I²C, bez obsługi RDS.

B.5. Kierunki rozbudowy i projekty pokrewne

Dla początkującego konstruktora przykładowy projekt jest dobrym punktem wyjścia w obecnej postaci – uczy obsługi I²C, SPI oraz cyfrowego nadajnika FM. Najprostsza modyfikacja to praca bez anteny zewnętrznej

Tabela 2. Porównanie scalonych nadajników FM

Układ	Typ/sterowanie	Stereo	RDS	Element odniesienia	Uwagi
KT0803L	cyfrowy, I ² C	tak (cyfrowe)	nie	rezonator 32,768 kHz	podstawa projektu z części A
QN8027	cyfrowy, I ² C	tak (cyfrowe)	nie	kwarc/rezonator	stabilny, ceniony za powtarzalność
Si4713	cyfrowy, I ² C	tak (cyfrowe)	tak	kwarc 32,768 kHz	najbogatszy funkcjonalnie, droższy
BH1417	PLL, przełącznik DIP	tak (PLL)	nie	kwarc 7,6/7,68 MHz	wyбір 14 częstotliwości, separacja ~40 dB
BA1404	analogowy, obwód LC	tak (analogowe)	nie	obwód LC (strojony)	klasyczny, separacja ~45 dB

(z krótkim odcinkiem przewodu jako anteną zastępczą), co utrzymuje moc w bezpiecznych granicach.

Dla konstruktora z doświadczeniem naturalnym rozwinięciem jest dodanie filtra dolnoprzepustowego na wyjściu antenowym oraz rozbudowa oprogramowania o regulację poziomu wejściowego i siły nadawania – funkcje, które KT0803L udostępnia, a których projekt w obecnej wersji nie wykorzystuje.

Dla zaawansowanego konstruktora interesującym kierunkiem jest przejście na układ Si4713 i dodanie transmisji RDS/RBDS – pozwala to wyświetlać nazwę stacji i komunikaty tekstowe w odbiorniku. Pokrewnym, dobrze udokumentowanym tematem jest cyfrowy odbiornik FM tego samego autora (na układzie TEA5767), który w naturalny sposób uzupełnia nadajnik w parę nadajnik–odbiornik.

Warto też sięgnąć po konstrukcje innych autorów, które pokazują alternatywne podejścia do tego samego zadania. Gotowy moduł Adafruit z układem Si4713 (projekt Limor Fried) to najprostsza droga do nadajnika z RDS – wraz z dopracowaną biblioteką Arduino i poradnikiem. Dla układu QN8027 dostępne są niezależne biblioteki sterujące z obsługą RDS (m.in. autorstwa M. Bhakara oraz M. Ondraka z Uniwersytetu w Pardubicach), a projekt VK6TT pokazuje sterowanie tym układem z mikrokontrolera STM8 w języku Forth – przykład realizacji poza ekosystemem AVR/Arduino. Odminną filozofię reprezentuje klasyczny nadajnik na układzie BH1417 (opracowanie EDN): częstotliwość wybiera się tam sprzętowo, przełącznikiem DIP, bez mikrokontrolera i magistrali cyfrowej. Pełne odnośniki zebrano poniżej.

Projekty pokrewne – odnośniki

Poniżej zebrano odnośniki do dokumentacji projektu przykładowego oraz konstrukcji pokrewnych. Adresy zweryfikowano pod kątem dostępności.

Projekt przykładowy (część A)

- [Hackster.io – Stereo Digital FM Transmitter Circuit](#) – strona projektu Hesama Moshiriego: opis, kod Arduino, schemat. [hackster.io/hesam-moshiri/stereo-digital-fm-transmitter-circuit-arduino-code-2dbd8d](#)
- [PCBWay – pełny artykuł i pliki Gerber](#) – rozszerzona wersja artykułu z kompletem ilustracji oraz plikami produkcyjnymi płytki. [pcbway.com/blog/technology/Stereo_Digital_FM_Transmitter_Circuit](#)
- [DigiKey Maker – Stereo Digital FM Transmitter Circuit](#) – ten sam projekt w bazie projektów DigiKey. [digikey.com/en/maker/projects/stereo-digital-fm-transmitter-circuit](#)

Projekty pokrewne tego samego autora

- [Hackster.io – Full Digital FM Receiver with Arduino and TEA5767](#) – cyfrowy odbiornik FM ze wzmacniaczem klasy D – uzupełnienie nadajnika w parę nadajnik–odbiornik. [hackster.io/hesam-moshiri/full-digital-fm-receiver-with-arduino-and-tea5767](#)
- [Hackster.io – FM Transmitter with RDS/RBDS Data Transmission](#) – nadajnik FM z transmisją danych RDS – kierunek rozbudowy projektu z części A. [hackster.io/hesam-moshiri/fm-transmitter-with-rds-rbds-data-transmission-capability](#)
- [Hackster.io – Digital Coil-Less FM Transmitter \(VMR6512\)](#) – prostszy nadajnik cyfrowy oparty na module VMR6512 – RF blok w jednej obudowie. [hackster.io/hesam-moshiri/how-to-build-a-digital-coil-less-fm-transmitter](#)

Projekty pokrewne innych autorów

- [Adafruit \(Limor Fried\) – Si4713 FM Transmitter Breakout](#) – kompletny moduł nadajnika stereo FM z RDS/RBDS oparty na układzie Si4713, wraz z biblioteką Arduino i poradnikiem – gotowy kierunek rozbudowy projektu o transmisję danych RDS. [learn.adafruit.com/adafruit-si4713-fm-radio-transmitter-with-rds-rbds-support](#)
- [GitHub – Adafruit-Si4713-Library](#) – biblioteka Arduino dla nadajnika Si4713/Si4714 z obsługą RDS (Adafruit, licencja BSD). [github.com/adafruit/Adafruit-Si4713-Library](#)
- [GitHub – ManojBhakarPCM/Arduino-QN8027-with-Full-RDS-support](#) – biblioteka Arduino dla układu QN8027 z pełną obsługą RDS i wszystkimi nastawami – alternatywa dla KT0803L w cyfrowym nadajniku FM. [github.com/ManojBhakarPCM/Arduino-QN8027-with-Full-RDS-support](#)
- [GitHub – dragon-engineer/QN8027 \(Martin Ondraka, Univerzita Pardubice\)](#) – biblioteka Arduino dla scalonego nadajnika FM z RDS QN8027 – projekt akademicki. [github.com/dragon-engineer/QN8027](#)
- [GitHub – VK6TT/QN8027-STM8-eForth](#) – sterowanie nadajnikiem QN8027 z mikrokontrolera STM8 w języku Forth – przykład realizacji na innej platformie niż AVR. [github.com/VK6TT/QN8027-STM8-eForth](#)
- [EDN – Stereo PLL FM Transmitter \(BH1417\)](#) – klasyczny stereofoniczny nadajnik FM z pętlą PLL na układzie BH1417, ze schematem i wykazem elementów – sprzętowy wybór częstotliwości przełącznikiem DIP. [edn.com/stereo-pll-fm-transmitter-bh1417](#)

Projekty i zestawy dostępne w polskich źródłach

Tematyka nadajników FM – zarówno scalonych, jak i dyskretnych – była wielokrotnie podejmowana w polskiej prasie elektronicznej.

Poniżej zebrano projekty i zestawy stanowiące odpowiedniki lub bliskie analogie konstrukcji omawianej w częściach A i B.

- [Elektronika Praktyczna – „transmitterFM. Miniaturowy nadajnik FM z RDS-em” \(EP 2/2014\)](#) – cyfrowy nadajnik FM na układzie Si4711/4713 z obsługą RDS, pasmo 87,5...108 MHz, sterowanie przez I²C – najbliższy polski odpowiednik projektu z części A, dodatkowo z transmisją danych RDS; dostępny w portalu EP.com.pl. [ep.com.pl/projekty/projekty-ep/9634-transmitterfm](#)
- [AVT5437 – zestaw „Miniaturowy nadajnik FM z RDS-em”](#) – kit do samodzielnego montażu odpowiadający projektowi z EP 2/2014 (Si4713, zasilanie 8...14 V DC, 4 komunikaty RDS typu PS) – dokumentacja w archiwum serwisowym AVT, zestaw bywa dostępny w sklepie AVT. [serwis.avt.pl/manuals/AVT5437.pdf](#)
- [Elektronika Praktyczna – „Nadajnik FM o mocy wyjściowej 2 W” \(EP 5/2000\)](#) – klasyczny nadajnik FM na elementach dyskretnych – generator W.cz. z osobnym stopniem mocy; pokazuje tradycyjne podejście, które układy scalone takie jak KT0803L wyparty. Materiał porównawczy do dyskusji w części B. [ep.com.pl/files/5419.pdf](#)
- [Sklep AVT – Radio FM z RDS, AVT5540](#) – Moduł odbiornika wykonano z użyciem popularnych i łatwych w montażu elementów, więc jego wykonanie nie sprawi trudności nawet początkującym. Prosty w użytkowaniu dzięki interfejsowi użytkownika zbudowanego z wyświetlacza LCD oraz dwóch impulsatorów. Zestaw posiada dodatkowy wzmacniacz o mocy 2×1 W

Dokumentacja układów i biblioteki

- [KT Micro – KT0803L](#) – nota katalogowa scalonego nadajnika FM zastosowanego w projekcie. [datasheetpdf.com – KT0803L](#)
- [GitHub – SSD1306Ascii \(greiman\)](#) – biblioteka obsługi wyświetlacza OLED SPI użyta w kodzie projektu. [github.com/greiman/SSD1306Ascii](#)
- [GitHub – MiniCore \(MCUdude\)](#) – menedżer płytek Arduino dla mikrokontrolerów ATmega8 i pokrewnych. [github.com/MCUdude/MiniCore](#)

Przepisy dotyczące nadajników małej mocy

- [EBU – raport R120](#) – stanowisko dotyczące nadajników FM małej mocy w Europie i limitu 50 nW ERP. [tech.ebu.ch/docs/r/r120.pdf](#)
- [ETSI EN 301357](#) – zharmonizowana norma dla bezprzewodowych urządzeń audio małej mocy (m.in. pasmo 87,5–108 MHz) – do pobrania w portalu normalizacyjnym ETSI. [etsi.org/standards \(norma EN 301357\)](#)

Your
B2B
partner

Tak! Zapobieganie przestojom w produkcji. Z Conrad.

Szybka dostawa pasujących części zamiennych



conrad.pl/tak-z-conrad

All parts of success

CONRAD

eprasa.pl 1761d60f5c

Nagrzewnice indukcyjne ZVS

Przykładowy projekt 1000 W oraz przegląd topologii konstrukcji od amatorskich do przemysłowych

Nagrzewanie indukcyjne od dawna fascynuje konstruktorów-amatorów: kilkadziesiąt watów wystarczy do efektownych pokazów, kilkaset – do wygrzewania połączeń i lutowania twardego, a kilka kilowatów pozwala hartować narzędzia i topić metale. Przegląd tej tematyki zawiera dwie części. W części A prezentujemy gotowy, sprawdzony projekt – nagrzewnicę indukcyjną 1000 W na rezonansowym obwodzie RLC autorstwa Borisa Landoni z serwisu Open-Electronics. W części B poddajemy ten projekt krytycznej dyskusji, wskazujemy miejsca warte poprawy, a następnie zestawiamy go z pokrewnymi rozwiązaniami – od najprostszych oscylatorów ZVS po przemysłowe mostki IGBT.

Część A. Przykładowy projekt Nagrzewnica indukcyjna 1000 W na rezonansowym obwodzie RLC

Stopmy metale techniką ZVS (Zero Voltage Switching – przełączanie przy zerowym napięciu) zastosowaną do rezonansowego obwodu RLC o mocy 1000 W. Czy zastanawialiście się kiedyś, jak silne potrafią być otaczające nas pola elektromagnetyczne?

Otacza nas niezliczona liczba fal elektromagnetycznych pochodzących z kabli linii elektroenergetycznych, nadajników radiowych i telewizyjnych, sieci telefonii komórkowej oraz pilotów zdalnego sterowania.

Ich oddziaływanie staje się zauważalne tylko w określonych sytuacjach i zależy od mocy promieniowania oraz naszej odległości od źródła. Ale jak silne może być pole elektromagnetyczne? Kuchenka mikrofalowa dowodzi, że potrafi ono podgrzać żywność, a nawet wywołać łuk elektryczny między brzegami folii aluminiowej lub metalowymi przedmiotami.

Czy takie pole może być na tyle silne, by w kilka sekund rozżarzyć metalowy przewód? Odpowiedź brzmi: tak – i prezentowany projekt znakomicie to pokazuje. To układ działający w oparciu o trzy zjawiska fizyczne: indukcję magnetyczną, prądy wirowe i efekt Joule'a.

Projekt pokazuje, jak nagrzewać, a nawet topić materiały przewodzące prąd – w szczególności ferromagnetyczne – za pomocą układu w technologii ZVS o mocy znamionowej 1000 W, która w pewnych warunkach może wzrosnąć do 1500 W. ZVS to technika

Autor projektu: Boris Landoni – Open-Electronics.org
Tytuł oryginału: „How to Build a 1000W ZVS Induction Heater Using a Resonant RLC Circuit”

stosowana w przekształtnikach energoelektronicznych w celu poprawy sprawności.

Pozwala ona przełącznikom półprzewodnikowym zmieniać stan przy napięciu na ich końcówkach bliskim zera, minimalizując straty mocy ($V \times I$), co wyjaśnimy w części poświęconej zasadzie działania. Ta koncepcja leży u podstaw indukcyjnych płyt kuchennych.

Indukcyjna płyta kuchenna składa się z cewki, przez którą płynie silny prąd przemienny, wytwarzający pole magnetyczne proporcjonalne do natężenia prądu. Zgodnie z prawem Faradaya zmienny w czasie strumień magnetyczny indukuje siłę elektromotoryczną (SEM) w każdym przewodzącym ciele, przecinanym przez linie wypadkowego pola.

Ta SEM generuje prądy wirowe wewnątrz materiału garnków lub naczyń postawionych na płycie. Prądy te nazywane są również prądami Foucaulta – ze względu na wirowy charakter ich przepływu wewnątrz przewodnika. Te prądy wywołują efekt Joule'a – wydzielanie ciepła wskutek strat energii. W indukcyjnych płytach kuchennych ciepło nagrzewa garnek.

Gotowanie indukcyjne jest przypadkiem szczególnym: efekt oddziaływania pól elektromagnetycznych zależy od przenikalności magnetycznej, reluktancji i przewodności użytych materiałów. To wyjaśnia, dlaczego metalowego garnka (zwłaszcza stalowego) nie wolno używać w kuchence mikrofalowej, a mięso można w niej

przyrządzać. I odwrotnie – na płycie indukcyjnej mięso się nie nagrzeje, jeśli nie znajduje się w naczyniu na bazie stali.

Zwykle prądy wirowe są niepożądane (jak choćby straty w rdzeniu transformatora), lecz w naszym przypadku celowo je wzmacniamy, aby wytworzyć ciepło w materiale poddanym działaniu pola magnetycznego. W szczególności materiały ferromagnetyczne nagrzewają się znacząco przy częstotliwości roboczej układu. Poza nagrzewaniem i topieniem metali układ ma też inne ciekawe zastosowania – bezprzewodowy przesył energii z użyciem sprzężonych cewek dostrojonych do rezonansu oraz układy zapłonowe oparte na cewkach Tesli.

Schemat ideowy nagrzewnicy

Analizując schemat ideowy (rysunek 1), w obwodzie sterującym dostrzeżemy symetryczną, dwugałęziową strukturę zwaną oscylatorem Royera, która zapewnia samowzbudne drgania bloku RLC na jego własnej częstotliwości rezonansowej. Blok RLC tworzą: cewka indukcyjna (zwana powszechnie cewką roboczą), bateria kondensatorów (nazywana kondensatorem obwodu rezonansowego) oraz rezystancja szeregową wnoszona przez elementy i połączenia.

Po włączeniu zasilania oscylator wchodzi w rezonans. Silny prąd rezonansowy płynący w cewce roboczej wytwarza intensywne pole magnetyczne. Po podaniu zasilania na oscylator Royera – nawet jeśli jest on symetryczny – jeden z dwóch tranzystorów MOSFET (M1 lub M2) zacznie przewodzić jako pierwszy, ze względu na rozrzut produkcyjny. Załóżmy, że jako pierwszy zacznie się otwierać M1 – jego dren zostaje ściągnięty do potencjału masy, więc M2 zostaje wyłączony przez diodę sprzężenia zwrotnego D2, która szybko rozładowuje bramkę M2.

Obwód rezonansowy złożony z cewki L3 i kondensatorów C1...C6 formuje na drenie M2 połówkę sinusoidy, narastającą od zera do wartości szczytowej i wracającą do zera. Gdy połówka sinusoidy wraca do zera, dioda D1 zaczyna przewodzić, wyłączając M1 i włączając M2, po czym

formuje się przeciwna połówka sinusoidy. Cykl ten powtarza się z częstotliwością równą własnej częstotliwości rezonansowej bloku RLC. Przez cewkę roboczą L3 płynie silny prąd sinusoidalny o częstotliwości około 100 kHz.

Tranzystory MOSFET pracują w układzie przeciwobnym (push-pull): gdy M1 jest włączony, M2 jest wyłączony – i na odwrót. Zapewniają to: obwód rezonansowy oraz diody sprzężenia zwrotnego D1 i D2. Jednoczesne przewodzenie spowodowałoby zwarcie źródła zasilania i zniszczenie tranzystorów.

Przełączanie MOSFET-ów następuje praktycznie przy zerowym napięciu dren-źródło (VDS), czyli realizowana jest technologia ZVS, co minimalizuje straty przełączania:

$$P_D = V_{DS} \cdot I_D$$

gdzie I_D to prąd drenu. ZVS znacząco ogranicza także zakłócenia wielkiej częstotliwości, naturalnie emitowane przez wszystkie układy przełączające.

Tranzystory M1 i M2 to Infineon **IRFP260N** – charakteryzują się bardzo małą rezystancją kanału w stanie otwarcia ($R_{DS(on)}$) zaledwie 40 mΩ oraz dużym prądem drenu ($I_D = 50$ A). Są to tranzystory szybkiego przełączania; dotyczy to również ich antyrównoległych diod ochronnych, których czas powrotu do stanu zaporowego (TRR) wynosi około 400 ns.

Napięcie szczytowe w układzie przełączającym ZVS oblicza się następująco:

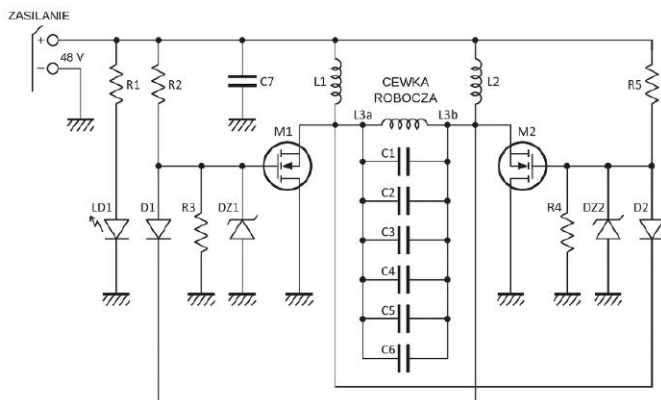
$$V_{PK} = V_{SUP} \cdot \pi$$

gdzie V_{SUP} to napięcie zasilania. Ponieważ napięcie zasilania jest ograniczone do 48 V DC, maksymalne napięcie dren-źródło (V_{DSS}) dobrano na poziomie 200 V.

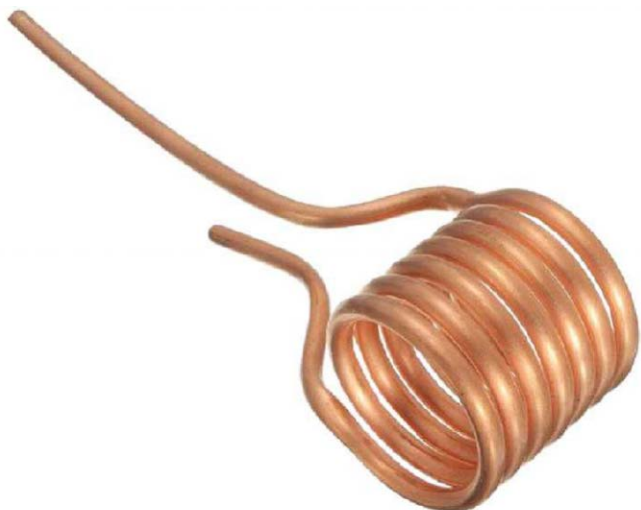
Diody sprzężenia zwrotnego D1 i D2 to MUR420 – również szybkie, zdolne przewodzić prąd co najmniej 4 A i szybko wyłączać odpowiedni MOSFET poprzez rozładowanie pojemności jego bramki. Bramki obu tranzystorów są sterowane przez rezystory mocy R2 i R5 o wartości 470 Ω każdy.

Aby chronić tranzystory MOSFET przed przepięciem i przeciążeniem prądowym, zastosowano dwie pary rezystor-dioda Zenera: R3/DZ1 oraz R4/DZ2. R3 i R4 mają 10 kΩ; DZ1 i DZ2 to diody 12 V/1 W. Dioda LED LD1 sygnalizuje obecność zasilania i jest zasilana przez rezystor szeregowy R1 (4,7 kΩ) ograniczający prąd.

Dławiki L1 i L2 mają indukcyjność 100 μH i służą do tłumienia impulsowych skoków napięcia, które mogłyby zniszczyć tranzystory MOSFET. Nawinięto je na rdzeniach toroidalnych i przewidziano na prąd maksymalny 13 A. Kondensatory C1–C6 to polipropylenowe kondensatory typu MKP, wybrane ze względu na zdolność przenoszenia dużych prądów i napięć podczas rezonansu przy minimalnych stratach. Pojemność rezonansowa wynosi



Rysunek 1. Schemat ideowy nagrzewnicy indukcyjnej – oscylator Royera/Mazzilliego z rezonansowym blokiem RLC



Rysunek 2. Cewka robocza złożona z rozsuniętych zwojów rurki miedzianej

w przybliżeniu $2 \mu\text{F}$ – uzyskano ją przez równoległe połączenie sześciu kondensatorów $0,33 \mu\text{F}/630 \text{ V}$.

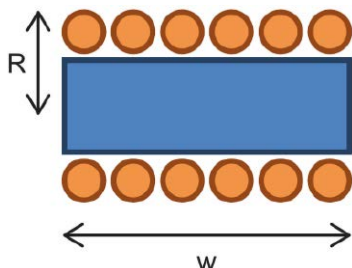
Cewka robocza

Cewka robocza (rysunek 2), do której wkłada się nagrzewany lub topiony obiekt, wykonana jest z 6 zwojów rurki miedzianej o średnicy 6 mm i grubości ścianki 0,8 mm. Jest nawinięta w powietrzu (bez rdzenia), co pozwala przenosić duże prądy, skutecznie odprowadzać wydzielane ciepło i ograniczać straty wynikające z efektu naskórkowości przy 100 kHz.

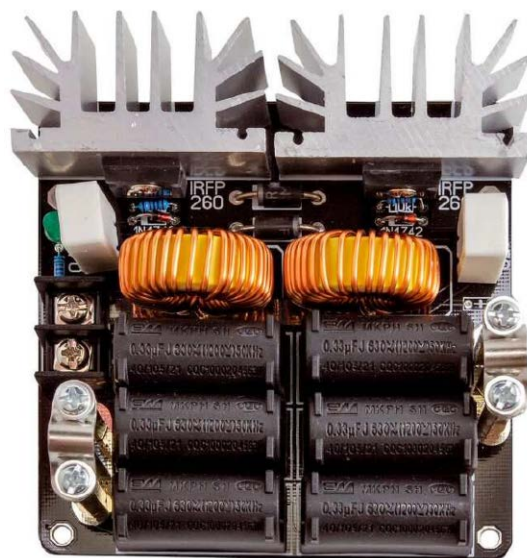
Zwoje nie mogą się stykać (miedź jest niezaizolowana), w przeciwnym razie powstaną zwarcia zmniejszające efektywną liczbę zwojów i zmieniające indukcyjność. Teoretyczna indukcyjność cewki roboczej wynosi około $1,26 \mu\text{H}$, choć może się zmieniać w zależności od długości poziomych wyprowadzeń łączących ją z kondensatorami. Do obliczenia teoretycznej indukcyjności własnej można posłużyć się wzorem (por. rysunek 3):

$$L = \frac{\mu_0 \cdot N^2 \cdot \pi \cdot R^2}{w}$$

gdzie μ_0 to przenikalność magnetyczna próżni ($\mu_0 = 4\pi \cdot 10^{-7} \text{ H/m}$), N to liczba zwojów (w naszym przypadku 6), $R = 2,5 \cdot 10^{-2} \text{ m}$, a $w = 7 \cdot 10^{-2} \text{ m}$; stąd $L \approx 1,26 \mu\text{H}$.



Rysunek 3. Cewka ma promień R równy 2,5 cm i długość w równą 7 cm



Rysunek 4. Zmontowana płytką nagrzewnicy

Teoretyczną częstotliwość rezonansową obliczamy ze wzoru:

$$f_R = \frac{1}{2\pi\sqrt{L \cdot C}} = 100,3 \text{ kHz}$$

gdzie $L = 1,26 \mu\text{H}$, $C = 2 \mu\text{F}$.

Montaż – schemat montażowy i wykaz elementów

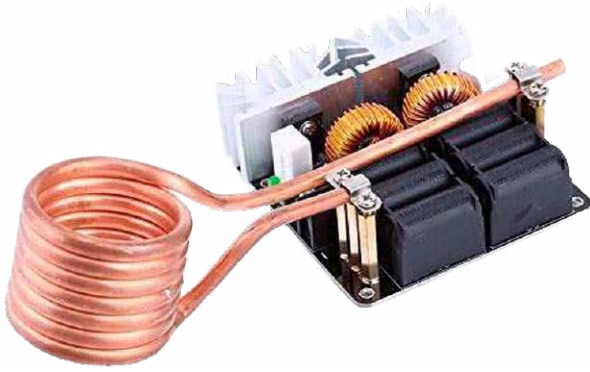
Zmontowaną płytkę pokazano na rysunku 4. Wykaz elementów zawiera tabela 1.

Montaż

Do zamocowania cewki użyto trzech sześciokątnych tulei dystansowych o średnicy 6 mm i długości 40 mm każda. Dwie boczne tuleje mocują (śrubami M4) stalowe obejmy unieruchamiające końce cewki roboczej. Należy poluzować śruby trzymające dwie tuleje (rysunek 5), wsunąć końce cewki roboczej pod obejmy tak, aby cewka była skierowana w stronę listwy zaciskowej zasilania oznaczonej „+/-”. Cewkę roboczą ustawia się tak, by

Tabela 1. Wykaz elementów

Oznaczenie	Element
C1...C6	$0,33 \mu\text{F}/630 \text{ V}$ -, raster 30 mm
C7	– (miejsce nieobsadzone)
R1	$4,7 \text{ k}\Omega$
R2	$470 \Omega/5 \text{ W}$
R3, R4	$10 \text{ k}\Omega/1\%$
R5	$470 \Omega/5 \text{ W}$
L1, L2	dławik $100 \mu\text{H}$
L3	cewka robocza $1,26 \mu\text{H}$
M1, M2	IRFP260N
D1, D2	MUR420
DZ1, DZ2	1N4742 (diody Zenera $12 \text{ V}/1 \text{ W}$)
LD1	diody LED zielona 5 mm



Rysunek 5. Tuleje dystansowe do mocowania cewki roboczej

spod obejmą wystawało mniej niż 1 cm rurki, po czym dokręca się śruby (**rysunek 6**).

Pomiary

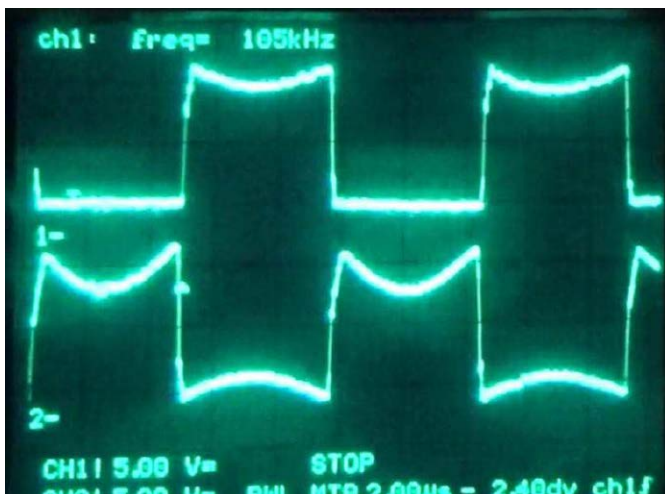
Wartość skuteczna prądu płynącego przez cewkę roboczą wynosi około 100 A. Przy tak dużym prądzie pole magnetyczne jest znaczne, a wszystkie znajdujące się w pobliżu przedmioty są wystawione na działanie jego linii sił – czułe urządzenia elektroniczne mogą więc doznawać zakłóceń. Gęstość strumienia magnetycznego B można obliczyć ze wzoru:

$$B = \frac{\mu_0 \cdot N \cdot I}{l} = 10,7 \text{ mT} = 107 \text{ Gs}$$

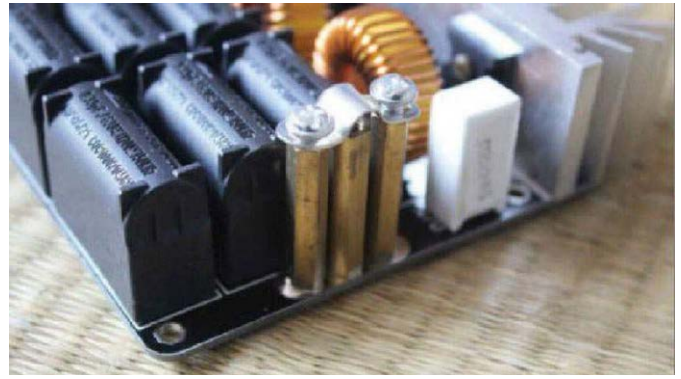
gdzie $N = 6$, $I = 100 \text{ A}$, $l = 7 \text{ cm}$ to długość cewki.

Na **rysunku 7** pokazano oscylogramy napięć bramka-źródło tranzystorów M1 i M2. Jak widać, są one idealnie w przeciwfazie: zanim napięcie na bramce M2 osiągnie poziom wysoki i M2 w pełni się otworzy, napięcie na bramce M1 jest już niskie – co gwarantuje, że M1 jest zatkany, i zapobiega jednoczesnemu przewodzeniu, które zniszczyłoby tranzystory.

Na **rysunku 8** pokazano napięcie na cewce roboczej – osiąga ono wartość szczytową rzędu 150 V i zachowuje



Rysunek 7. Oscylogramy napięć bramka-źródło tranzystorów M1 i M2 – przebiegi w przeciwfazie



Rysunek 6. Cewka robocza zamontowana na płycie

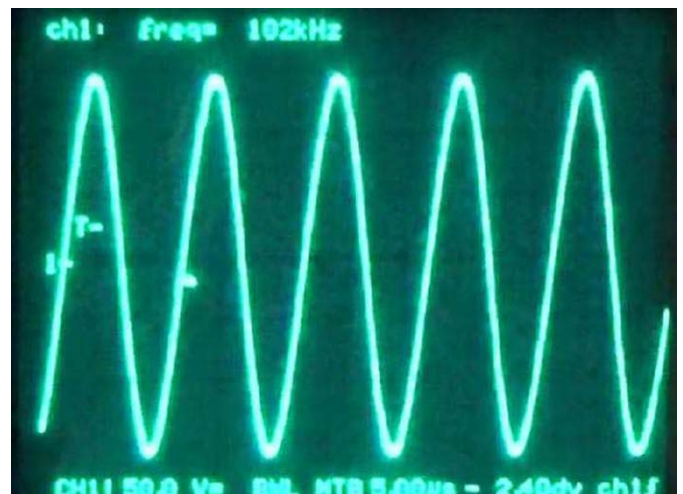
idealnie sinusoidalny przebieg dzięki równoległemu rezonansowemu obwodowi RLC.

Na koniec **rysunek 9** przedstawia prąd pobierany przez układ przy napięciu wejściowym 48 V DC, wynoszący około 17 A. Prąd zmierzono, odczytując spadek napięcia na boczniku 0,06 Ω (uzyskanym przez równoległe połączenie trzech rezystorów 0,18 Ω) za pomocą multimetru cyfrowego, z metalowym przedmiotem umieszczonym w cewce roboczej.

Uruchomienie i użytkowanie

Układ należy zasilać napięciem stałym z zakresu od 10 do 48 V o wystarczającej mocy. W przypadku stosowania maksymalnego napięcia 48 V DC zalecany jest zasilacz o mocy co najmniej 1500 W. Z zachowaniem biegunowości wsuwa się przewody dodatni i ujemny zasilacza do listwy zaciskowej i dokręca śruby. Po podaniu zasilania należy sprawdzić, czy zaświeciła się zielona dioda LED LD1 sygnalizująca obecność napięcia wejściowego.

Od tej chwili przez cewkę roboczą zaczyna płynąć prąd rezonansowy o częstotliwości około 100 kHz. Przy zasilaniu pełnym napięciem 48 V i braku jakichkolwiek metalowych przedmiotów wewnątrz cewki pobór mocy nie przekracza 500 W; po włożeniu przewodnika – w zależności



Rysunek 8. Oscylogram napięcia na cewce roboczej – przebieg sinusoidalny o wartości szczytowej ok. 150 V



Rysunek 9. Pomiar prądu pobieranego przez układ na podstawie spadku napięcia na boczniku



Rysunek 10. Obiekt umieszczony w cewce roboczej rozżarza się w ciągu kilku sekund

od jego materiału, rozmiaru, kształtu i położenia – pobór mocy może się jednak silnie zmieniać i sięgać 1500 W.

Przewodnik, który chcemy umieścić w cewce – na przykład metalową śrubę – należy trzymać szczypcami (nie zbliżając się zbyt do cewki), aby uniknąć poparzeń. Przewodnik wsuwa się powoli do cewki roboczej, uważając, by nie dotykał jej wewnętrznych ścianek, i utrzymuje w pozycji pionowej.

Wsunięty przewodnik działa jak uzwojenie wtórne transformatora, a cewka robocza – jak uzwojenie pierwotne. Znajdując się wewnątrz solenoidu, przewodnik jest poddany działaniu silnego pola magnetycznego. W ciągu kilku sekund, dzięki efektowi Joule'a, przewodnik nagrzewa się i zaczyna świecić charakterystycznymi pomarańczowo-żółtymi barwami żarzenia (rysunek 10).

Stan ten można utrzymywać przez kilkadziesiąt sekund. Należy jednak zachować szczególną ostrożność, ponieważ temperatura szybko rośnie – aby uniknąć poparzeń, nie wolno dotykać bezpośrednio przewodnika, cewki ani radiatorów. Po kilku minutach pracy cewka robocza zacznie ciemnieć i z czasem stanie się niemal czarna od wydzielanego ciepła. Nie wpływa to na jej parametry.

Część B. Dyskusja i projekty pokrewne

Część B składa się z dwóch wątków. Najpierw poddajemy przykładowy projekt krytycznej ocenie – wskazujemy miejsca, w których oryginalny opis bywa nieścisły, oraz rozwiązania warte poprawy. Następnie umieszczamy konstrukcję Borisa Landoni w szerszym kontekście: przeglądamy rodziny topologii nagrzewnic indukcyjnych – od najprostszych oscylatorów ZVS po przemysłowe mostki IGBT – wraz z zestawieniami porównawczymi, doborem komponentów i typowymi problemami.

B.1. Komentarze do przykładowego projektu

Projekt Borisa Landoni to dojrzała, wielokrotnie powielana konstrukcja należąca do najpopularniejszej rodziny

nagrzewnic amatorskich – oscylatora ZVS typu Royer/Mazzilli. Jest atrakcyjny prostotą i powtarzalnością; warto jednak zwrócić uwagę na kilka miejsc, w których oryginalny opis bywa nieścisły lub w których konstrukcję dałoby się ulepszyć.

1. Margines napięciowy tranzystorów jest niewielki.

Autor poprawnie podaje, że napięcie szczytowe na drenie wynosi $V_{PK} = \pi \cdot V_{SUP}$. Przy 48 V daje to $\pi \cdot 48 \approx 151$ V. Tranzystor IRFP260N ma napięcie $V_{DS} = 200$ V, zatem zapas wynosi zaledwie ok. 50 V (ok. 25%). W rzeczywistym układzie na drenach pojawiają się dodatkowe oscylacje i przebiegi ponad wartość idealną πV_{CC} – wynikające z indukcyjności rozproszenia połączeń i dławików. To jeden z głównych powodów, dla których egzemplarze pracujące przy pełnych 48 V bywają uszkodzane. W praktyce 48 V należy traktować jako absolutny kres możliwości IRFP260N; dla zachowania marginesu bezpieczeństwa rozsądniej pracować z zasilaniem do 36...40 V, a przy zasilaniu 48 V – zadbać o bardzo krótkie, „grube” połączenia drenów z baterią kondensatorów oraz nie liczyć na długotrwałą pracę przy mocy 1500 W.

2. Diody Zenera DZ1/DZ2 chronią bramki, a nie dreny.

Oryginał podaje, że pary R3/DZ1 i R4/DZ2 „chronią tranzystory MOSFET przed przebiegiem i przeciążeniem prądowym”. To opis nieścisły. Diody Zenera 12 V włączone między bramkę a źródło są ogranicznikami napięcia bramki: w oscylatorze Royera bramka jest „ciągnięta” w górę przez rezystor i sprzężona z przeciwnym drenem, więc bez ogranicznika napięcie bramka-źródło mogłoby przekroczyć dopuszczalne ± 20 V tranzystora IRFP260N i przebić cienką warstwę tlenku bramki. DZ1/DZ2 utrzymują VGS na bezpiecznym poziomie ok. 12 V. Nie pełnią natomiast funkcji ochrony drenu przed przebiegiem ani ograniczania prądu – takich funkcji ta topologia po prostu nie ma.

3. Czas TRR 400 ns dotyczy diod podłoża tranzystorów.

Tekst podaje, że tranzystory są szybkie, „co dotyczy również ich antyrównoległych diod ochronnych

o TRR ok. 400 ns”. Warto doprecyzować: chodzi o diody podłoża (body diode) tranzystorów IRFP260N, które w układzie przeciwobnym ZVS przewodzą w części cyklu. Ich stosunkowo długi czas powrotu do stanu zaporowego (rzędu setek ns) jest źródłem strat przełączania i zakłóceń – to znana słaba strona tej rodziny tranzystorów. Diody sprzężenia zwrotnego D1/D2 (MUR420) są natomiast diodami ultraszybkimi i do swojej funkcji – błyskawicznego rozładowania bramki – w zupełności wystarczają.

4. **Dławiki zasilające L1/L2 są na granicy obciążalności prądowej.** Wykaz elementów podaje dławiki 100 μ H o prądzie maksymalnym 13 A. Tymczasem przy mocy wejściowej dochodzącej do 1500 W i napięciu 48 V układ pobiera nawet ok. 31 A. Każdy z dwóch dławików przenosi składową stałą rzędu połowy tego prądu (ok. 15 A) – a więc powyżej deklarowanych 13 A. Przy pełnej mocy rdzeń może wchodzić w nasycenie, indukcyjność spada, a odsprzęganie zasilacza od obwodu rezonansowego się pogarsza. Zaleca się dobór dławików o zapasie indukcyjności i prądu nasycenia odpowiednim do rzeczywistego prądu wejściowego – albo świadome ograniczenie mocy/napięcia pracy.
5. **Brak miękkiego startu.** Topologia Royera/Mazzilliego nie ma układu łagodnego rozruchu – musi być załączana skokowo, od 0 V do pełnego napięcia. Ma to dwie praktyczne konsekwencje. Po pierwsze, zasilacz laboratoryjny z ograniczeniem prądowym może uniemożliwić wzbudzenie drgań – oscylator potrzebuje pełnego skoku napięcia. Po drugie, przy zasilaczu impulsowym skok prądu rozruchowego (ładowanie baterii kondensatorów i wejście w rezonans) może wprowadzić zasilacz w stan zabezpieczenia. Warto przewidzieć obwód miękkiego startu – np. rezystor ograniczający, zwierany po chwili przekaźnikiem – lub dobrać zasilacz tolerujący taki rozruch.
6. **Kondensator C7 pozostaje nieobsadzony.** Na schemacie i płytce przewidziano kondensator C7 na wejściu zasilania, ale w wykazie elementów ma on wartość „–” (nie montowany). Pozostawienie tego miejsca pustym oznacza brak odsprzęgania linii zasilającej. Długie przewody zasilacza tworzą wraz z indukcyjnościami obwodu pętlę zdolną do oscylacji i podbijania przepięć na drenach. Zaleca się obsadzenie C7 – najlepiej kondensatorem foliowym o niskim ESR (rzędu 1...10 μ F), ewentualnie w połączeniu z kondensatorem elektrolitycznym dużej pojemności jako zasobnikiem energii. Poprawia to stabilność i zmniejsza obciążenie tranzystorów.
7. **Cewka robocza pracuje „na sucho”.** Cewkę wykonano z rurki miedzianej – a więc elementu wprost

przeznaczonego do chłodzenia cieczą. W opisywanej konstrukcji rurka nie jest jednak przepłukiwana wodą; autor sam zaznacza, że cewka „ciemnieje i staje się niemal czarna” od ciepła. To jasny sygnał ograniczonego cyklu pracy. Dla pracy ciągłej przy mocy zbliżonej do znamionowej rurkę warto włączyć w obieg wody – wtedy spełnia ona swoją właściwą rolę. Bez chłodzenia nagrzewnicę należy traktować jako urządzenie o pracy dorywczej (kilkadziesiąt sekund, jak pisze autor).

8. **Brak zabezpieczenia wejścia.** Zalecany zasilacz 1500 W/48 V to źródło o bardzo dużej energii – prąd zwarciový jest groźny zarówno dla układu, jak i dla otoczenia. W projekcie nie przewidziano bezpiecznika na wejściu DC. Zaleca się dodanie bezpiecznika dobrego do prądu roboczego (lub poleganie na sprawnym zabezpieczeniu nadprądowym zasilacza) oraz prowadzenie przewodów zasilających możliwie krótkich i o odpowiednim przekroju.
9. **Moc znamionowa a moc rzeczywista.** Tytułowe „1000 W” jest wartością raczej hasłową. Zmierzony pobór 17 A przy 48 V to ok. 816 W mocy wejściowej; osiągnięcie 1000 W wymaga ok. 21 A, a 1500 W – ok. 31 A (stąd zalecenie zasilacza 1500 W). Przy chłodzeniu powietrznym i tranzystorach IRFP260N praca przy najwyższych mocach jest możliwa wyłącznie dorywczo – ograniczeniem jest nagrzewanie tranzystorów i baterii kondensatorów. To realistyczne ujęcie, spójne z doświadczeniami społeczności budującej podobne moduły.
10. **Częstotliwość pracy a rodzaj wsadu.** Podana częstotliwość $f_R \approx 100$ kHz to wartość teoretyczna, dla cewki bez wsadu. Po włożeniu obiektu – zwłaszcza ferromagnetycznego – efektywna indukcyjność i sprzężenie się zmieniają, a częstotliwość pracy maleje. Jest to normalne i pożądane: oscylator ZVS samoczynnie „śledzi” rezonans i nie wymaga układu regulacji. Warto jednak pamiętać, że do głębokiego nagrzewania większych wsadów stalowych korzystniejsze są niższe częstotliwości (rzędu kilkudziesięciu kHz), uzyskiwane przez zwiększenie pojemności baterii lub liczby zwojów cewki. Dla małych elementów 100 kHz jest w pełni odpowiednie.
11. **Liczba zwojów cewki – rozbieżność z dokumentacją oryginału.** Uważny czytelnik dostrzeże nieścisłość między tekstem a fotografiami. W opisie cewki roboczej (część A) oraz we wszystkich wzorach autor przyjmuje sześć zwojów, tymczasem na zdjęciach gotowego urządzenia – zwłaszcza na rysunkach 2 i 6 – wyraźnie widać, że nawinięta cewka ma ich siedem. Fotografia jest tu dowodem rozstrzygającym: rzeczywista cewka prezentowanego egzemplarza jest

siedmiozwojowa. Rozbieżność warto skomentować, ponieważ liczba zwojów wprost wpływa na obliczenia. Po pierwsze, indukcyjność. Uproszczony wzór geometryczny podany w części A daje dla $N = 6$ wartość $L \approx 1,27 \mu\text{H}$, a dla $N = 7$ – już $L \approx 1,73 \mu\text{H}$ (indukcyjność rośnie z kwadratem liczby zwojów). Po drugie, częstotliwość rezonansowa. Podstawiając $L \approx 1,73 \mu\text{H}$ do wzoru na f_R , otrzymalibyśmy około 86 kHz zamiast podanych 100,3 kHz. Tymczasem oscylogramy w części A („Pomiary”) pokazują częstotliwość rzędu 100...105 kHz – a więc wartość zgodną z obliczeniem dla sześciu zwojów, nie siedmiu. Po trzecie, indukcja magnetyczna. We wzorze na B liczba zwojów występuje w pierwszej potęgze, więc dla $N = 7$ wynik rośnie proporcjonalnie: $B \approx 12,6 \text{ mT}$ (126 Gs) zamiast 10,7 mT. Jak to pogodzić? Najbardziej prawdopodobne wyjaśnienie jest takie, że autor oryginału przyjął w obliczeniach $N = 6$, podczas gdy zbudowany model ma siedem zwojów – przy czym podane $L \approx 1,26 \mu\text{H}$ i $f_R \approx 100 \text{ kHz}$ najlepiej traktować jako wartości zmierzone (potwierdzone oscylogramami), a sam wzór geometryczny – jako zgrubne oszacowanie. Wzór ten pomija bowiem odstęp między zwojami, skończoną średnicę rurki oraz indukcyjność poziomych wyprowadzeń, więc jego zgodność z pomiarem akurat dla $N = 6$ jest po części przypadkowa. Dla konstruktora powtarzającego projekt wniosek jest praktyczny: liczba zwojów nie jest wartością krytyczną – oscylator ZVS samoczynnie dostraja się do rezonansu (por. komentarz 10), a ostateczna częstotliwość i tak

zależy od pojemności baterii kondensatorów oraz od wsadu. Cewkę można więc nawinąć z sześciu lub siedmiu zwojów; należy jedynie mieć świadomość, że podane we wzorach liczby odnoszą się do wariantu sześciuzwojowego.

Podsumowanie. Mimo powyższych uwag projekt zostaje znakomitą punktem wyjścia dla konstruktora-amatora – jest prosty, powtarzalny i efektowny. Wskazane poprawki (margines napięciowy, dobór dławików, miękki start, odsprężanie wejścia, chłodzenie cewki, zabezpieczenie nadprądowe) podnoszą niezawodność i pozwalają bezpiecznie zbliżyć się do deklarowanej mocy.

B.2. Topologie nagrzewnic indukcyjnych – przegląd

Każda nagrzewnica indukcyjna wytwarza w cewce roboczej szybkozmienne pole magnetyczne, które indukuje prądy wirowe w umieszczonym wewnątrz metalu; efekt Joule’a nagrzewa wsad bez kontaktu fizycznego. Różnice między konstrukcjami sprowadzają się głównie do sposobu sterowania kluczami półprzewodnikowymi oraz do poziomu mocy. Technika ZVS – przełączanie przy zerowym napięciu – jest wspólnym mianownikiem większości rozwiązań amatorskich, ponieważ eliminuje straty przełączania i ogranicza emisję zakłóceń.

Oscylator Royera/Mazzilliego (ZVS push-pull)

Najpopularniejsza i najprostsza topologia amatorska – to właśnie rodzina, do której należy przykładowy projekt

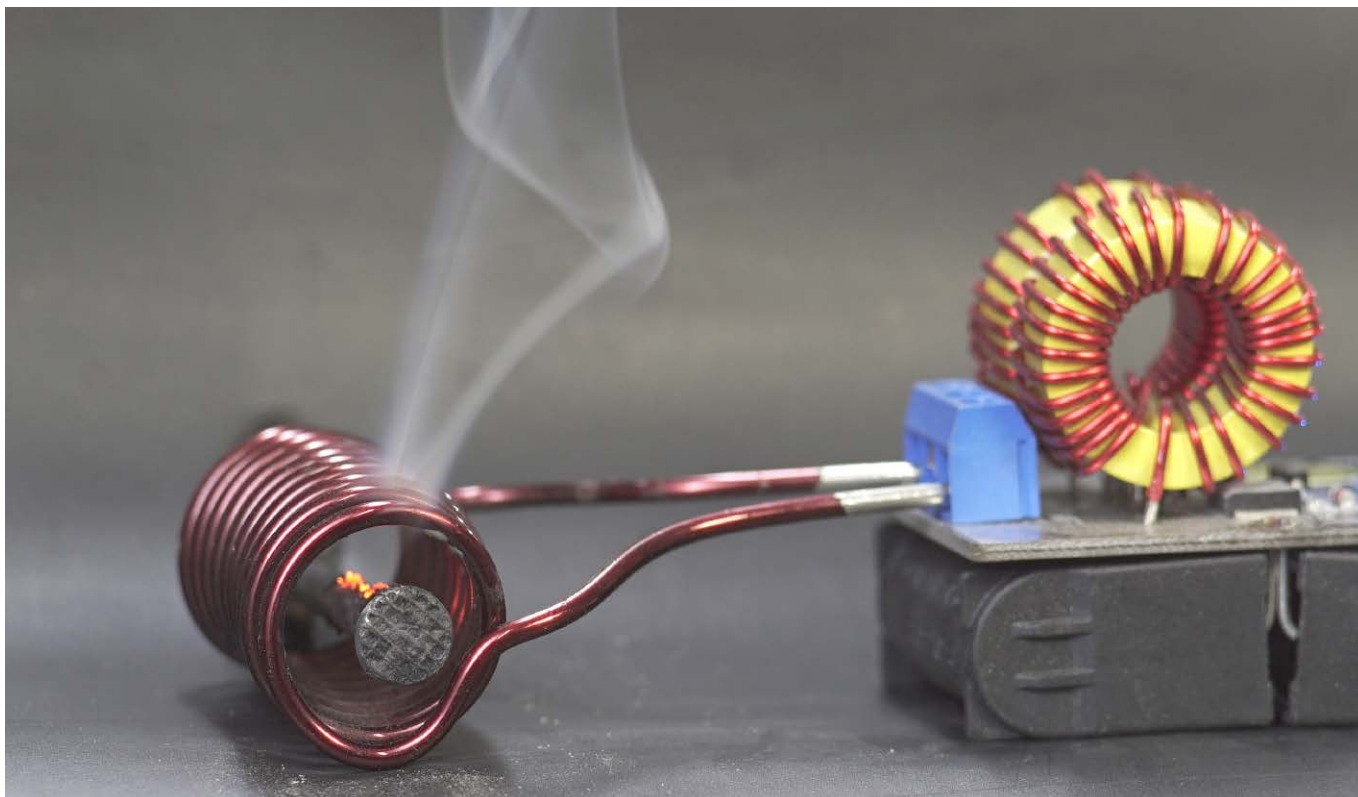


Tabela 2. Porównanie projektów nagrzewnic indukcyjnych DIY

Projekt/topologia	Zasilanie	Moc wej.	Częstotliwość	Elementy czynne	Trudność	Koszt części
Klasyczny ZVS Mazzilli	12...36 V DC	200...500 W	50...150 kHz (autorezonans)	2× IRFP250N/IRFP260N	★★★★ łatwy	15...35 €
ZVS dużej mocy	24...48 V DC	1...2 kW	80...130 kHz	4...8× IRFP260N/IRFP4668	★★★★ średni	50...120 €
ZVS Royer z chłodzeniem wodnym	12...48 V DC	do 1,4 kW	80...120 kHz	2× IRFP260N/IRFP4668 + chłodzenie	★★★★ średni	60...100 €
Mostek półtłokowy IR2153 + IGBT	60...400 V DC	0,5...2 kW	12...100 kHz (PFM)	2× STGW30NC60W/IRG4PC50S	★★★★ zaawansowany	80...150 €
Mostek pełny IGBT	200...400 V DC	3...12 kW	20...80 kHz	4× FGA25N120/CM300DY-24A	★★★★ ekspert	300...1500 €
Quasi-rezonansowy (kuchenny)	220...230 V AC	do 2 kW	20...50 kHz	1× IGBT 600 V + dioda	★★★★ zaawansowany	30...60 €
Royer LiPo/mobilny	11...25 V (LiPo)	100...500 W	100...300 kHz	2× IRF540N/IRFZ44N	★★★★ łatwy	10...25 €

z części A. Bywa nazywana „klasycznym ZVS” lub „driverem Mazzilliego”. Jest to oscylator przeciwobny z rezonansową cewką roboczą, w którym dreny obu tranzystorów MOSFET są krzyżowo sprzężone z bramkami tranzystora przeciwnego. Częstotliwość drgań ustala się samoczynnie, zgodnie z parametrami obwodu LC, według zależności $f = 1/(2\pi\sqrt{LC})$. Napięcie na drenach przebiega sinusoidalnie i może osiągać πV_{CC} , dlatego tranzystory dobiera się na napięcie rzędu czterokrotności napięcia zasilania. Wariant dużej mocy stosuje po kilka tranzystorów w każdej gałęzi, kondensatory MKP wysokiej klasy i chłodzenie wodne cewki roboczej.

Mostek półtłokowy z modulacją częstotliwości

Topologia przemysłowa, adaptowana do zaawansowanych projektów amatorskich. Stosuje dedykowany układ sterujący (np. IR2153, IRS2153, EG3005) do generowania sygnałów bramkowych o regulowanej częstotliwości.

Szeregowy obwód rezonansowy LC pracuje powyżej częstotliwości rezonansowej, co zapewnia warunek ZVS; moc reguluje się modulacją częstotliwości (PFM). Zasilanie pochodzi zwykle z prostownika sieciowego (60...400 V DC).

Mostek pełny IGBT

Topologia przemysłowa do dużych mocy, powyżej 3 kW. Cztery tranzystory IGBT w układzie mostka pełnego, sterowane przez procesor sygnałowy DSP lub dedykowany kontroler rezonansowy LLC. Stosowana w piecach do topienia metali i przemysłowych hartowniach. Wymaga wiedzy z zakresu energoelektroniki oraz rygorystycznego podejścia do bezpieczeństwa elektrycznego.

Topologia quasi-rezonansowa (jeden IGBT)

Najprostsza topologia zasilana wprost z sieci 230 V AC – jeden tranzystor IGBT z antyrównoległą diodą, regulacja przez PFM. To rozwiązanie stosowane w domowych

Tabela 3. Tranzystory mocy stosowane w nagrzewnicach indukcyjnych

Symbol	Typ	V_{DS}/V_{CE}	I_D/I_C	RDS(on)	Zastosowanie	Cena (1 szt.)
IRFP250N	MOSFET N	200 V	30 A	85 mΩ	ZVS do 36 V/500 W	ok. 1...2 €
IRFP260N	MOSFET N	200 V	50 A	40 mΩ	ZVS do 48 V/1 kW (preferowany)	ok. 2...3 €
IRFP4668	MOSFET N	200 V	130 A	10 mΩ	ZVS dużej mocy, małe RDS(on)	ok. 4...6 €
IRF540N	MOSFET N	100 V	33 A	44 mΩ	ZVS do 24 V/300 W (LiPo)	ok. 0,5...1 €
STGW30NC60W	IGBT	600 V	30 A	$V_{CE} = 1,5 V$	Mostek półtłokowy 230 V AC/1 kW	ok. 2...4 €
IRG4PC50S	IGBT	600 V	55 A	$V_{CE} = 1,8 V$	Mostek półtłokowy 230 V AC/2 kW	ok. 4...6 €
FGA25N120ANTD	IGBT	1200 V	25 A	–	Mostek pełny 400 V/3 kW	ok. 5...8 €

Tabela 4. Typy kondensatorów w obwodach rezonansowych

Rodzaj	Dielektryk	Pojemność	Napięcie	Uwagi
MKP	polipropylen (metalizowany)	0,1...1 μF/szt.	630...1600 V DC	Wymagane – niskie ESR, niskie straty przy 100 kHz
KP/KP-SN	polipropylen (zwój foliowy)	ok. 0,33 μF	1000...1600 V	Idealne do ZVS – np. wyroby RIFA/EPCOS
MKP klasy X2	polipropylen bezpieczny	0,1...0,47 μF	275...305 V AC	Do mostka półtłokowego; nie do ZVS powyżej 36 V
Mikowy	mika	0,01...0,1 μF	500...2000 V	Bardzo niskie straty, drogi, do ok. 200 kHz
Elektrolityczny	–	–	–	Bezwzględnie zakazany w obwodzie rezonansowym

Tabela 5. Wymagania zasilania według topologii

Topologia	U min.	U maks.	Prąd typowy	Uwagi
ZVS Mazzilli (mały)	10 V DC	36 V DC	5...20 A	Akumulator 12/24 V lub zasilacz laboratoryjny
ZVS dużej mocy	24 V DC	48 V DC	20...50 A	Zasilacz impulsowy 24 V/40 A lub 48 V/30 A
Mostek połówkowy IR2153	60 V DC	400 V DC	5...15 A	Prostownik 230 V AC + kondensatory 400 V
Mostek pełny IGBT	200 V DC	400 V DC	15...40 A	Zasilanie trójfazowe lub układ PFC
Quasi-rezonansowy	220 V AC	230 V AC	5...10 A	Bezpośrednio z sieci – brak izolacji
Royer LiPo (mobilny)	11 V	25 V	10...30 A	Pakiet LiPo 3S...6S; konieczny układ BMS

Tabela 6. Typowe problemy w nagrzewnicach ZVS

Problem	Przyczyna	Rozwiązanie
Uszkodzenie MOSFET-ów przy starcie	Brak wzbudzenia drgań przy zbyt niskiej dobroci Q obwodu	Sprawdzić wartość kondensatorów MKP; nie uruchamiać bez cewki roboczej
Brak drgań przy silnym obciążeniu	Zbyt silne sprzężenie z metalem (zbyt małe Q)	Zwiększyć liczbę zwojów cewki lub szczelinę między cewką a wsadem
Przegrzewanie MOSFET-ów	Zbyt duże RDS(on) lub zbyt małe radiatory	Zastosować IRFP260N lub IRFP4668; użyć podkładek izolacyjnych Al_2O_3
Nagrzewanie kondensatorów	Niewłaściwy typ lub zbyt duży prąd na jeden kondensator	Stosować wyłącznie MKP/KP; rozdzielić prąd na baterię kondensatorów
Wysoki poziom zakłóceń (EMI)	Brak filtra na wejściu zasilania	Dodać filtr LC na wejściu; ekranować obwód rezonansowy
Niestabilna częstotliwość pracy	Zmiana indukcyjności cewki przy wkładaniu wsadu	Zjawisko normalne – topologia ZVS śledzi rezonans samoczynnie

plytach indukcyjnych, o mocy do ok. 2 kW. Dla konstrukcji amatorskich jest niezalecane ze względu na brak izolacji galwanicznej od sieci.

Topologia LCLR z transformatorem dopasującym

Wariant zaawansowany: szeregowy dławik dopasowujący w połączeniu z równoległym obwodem rezonansowym cewki. Pozwala pracować przy niskim napięciu zasilania (24...48 V DC), przetwarzając je – poprzez transformator ferrytowy – na wysokie prądy cewki. Stosowany w projektach o mocy od 1 do kilkunastu kilowatów.

B.3. Zestawienie porównawcze projektów

Tabela 2 zestawia typowe projekty amatorskie według topologii i poziomu mocy. Przykładowy projekt z części A mieści się pomiędzy wierszami „ZVS dużej mocy 1...2 kW” a „ZVS Royer z chłodzeniem wodnym” – jest klasycznym, jednoparowym oscylatorem ZVS skalowanym do ok. 1 kW.

B.4. Kluczowe komponenty Tranzystory MOSFET i IGBT

Podstawowa zasada doboru: napięcie VDSS tranzystora MOSFET musi wynosić co najmniej czterokrotność napięcia zasilania, ponieważ rezonansowy wzrost napięcia w obwodzie LC powoduje, że napięcie w obwodzie rezonansowym sięga πV_{CC} .

Tabela 3 zestawia tranzystory najczęściej stosowane w nagrzewnicach amatorskich.

Kondensatory rezonansowe

Kondensatory obwodu rezonansowego pracują w najtrudniejszych warunkach całego układu – przenoszą duże prądy o częstotliwości rzędu 100 kHz.

Wymagane są typy o niskim ESR i niskich stratach; w praktyce oznacza to kondensatory polipropylenowe MKP lub KP.

W przykładowym projekcie zastosowano sześć kondensatorów 0,33 μF /630 V połączonych równolegle (łącznie około 2 μF) – ich montaż na szynie miedzianej zapewnia równomierny rozkład prądu.

Kondensatorów elektrolitycznych nie wolno stosować w obwodzie rezonansowym pod żadnym pozorem.

Dławiki zasilające

Dławiki L1/L2 w topologii ZVS pełnią dwie funkcje: izolują zasilacz od rezonującego obwodu LC oraz zapobiegają „zalaniu” rezonansu impedancją zasilacza. Typowe wartości to 45...200 μH , a – co kluczowe i o czym była mowa w komentarzu nr 4 – prąd nasycenia dławika musi przewyższać rzeczywisty prąd zasilania. Najczęściej stosuje się rdzenie toroidalne ferrytowe lub proszkowe z 8...15 zwojami drutu nawojowego o przekroju 1,5...2,5 mm².

B.5. Wymagania dotyczące zasilania

Dobór zasilacza jest równie istotny jak dobór samej nagrzewnicy. Tabela 5 podaje typowe zakresy napięć i prądów dla poszczególnych topologii. Warto pamiętać, że oscylator ZVS – jak omówiono w komentarzu nr 5 – należy załączać skokowo pełnym

napęciem; zasilacze z agresywnym ograniczeniem prądowym mogą uniemożliwić wzbudzenie drgań.

B.6. Typowe problemy i rozwiązania

Większość usterek nagrzewnic ZVS sprowadza się do kilku powtarzalnych przyczyn. Tabela 6 zbiera najczęstsze z nich wraz z zalecanym postępowaniem.

B.7. Zalecenia przy wyborze projektu

Dla początkującego konstruktora-amatora najlepszym wyborem jest klasyczny oscylator ZVS Mazzilli o mocy ok. 500 W z tranzystorami IRFP260N, czterema kondensatorami MKP0,33 μ F/630 V i zasilaczem 24 V/20 A. To dokładnie rodzina przykładowego projektu z części A, tyle że w mniejszej skali – prosta, tania i wybacza-jąca błędy.

Dla konstruktora z doświadczeniem (hartowanie, topienie metali nieżelaznych) odpowiedni jest oscylator ZVS o mocy 1...2 kW z chłodzeniem wodnym cewki i zasilaczem 48 V/30 A – czyli przykładowy projekt rozbudowany zgodnie z komentarzami z punktu B.1: z poprawnie dobranymi dławikami, obwodem miękkiego startu, obsadzonym kondensatorem wejściowym i przepływowym chłodzeniem cewki.

Dla zaawansowanego konstruktora (zastosowania zbliżone do przemysłowych) właściwym kierunkiem jest mostek półkowy z układem IR2153 i tranzystorami IGBT, zasilany z sieci 230 V AC przez prostownik. Topologia ta daje regulację mocy i lepiej skaluje się w górę, wymaga jednak izolowanych sterowników bramek i rygorystycznego podejścia do bezpieczeństwa pracy z napięciem sieciowym.

Projekty pokrewne – odnośniki

Poniżej zebrano sprawdzone odnośniki do dokumentacji projektów omawianych w częściach A i B. Adresy zweryfikowano pod kątem dostępności; przy projektach udostępniających pliki produkcyjne zaznaczono ich zakres (schemat, pliki Gerber, projekt KiCad).

Projekt przykładowy (część A)

- [Open-Electronics – 1000 W ZVS Induction Heater](https://open-electronics.org/how-to-build-a-1000w-zvs-induction-heater) – oryginał omawianego artykułu Borisa Landoni: schemat, wykaz elementów, pomiary. open-electronics.org/how-to-build-a-1000w-zvs-induction-heater
- [Open-Electronics – An Open Project for a 1000 W Induction Heater](https://open-electronics.org/an-open-project-for-a-1000w-induction-heater) – wcześniejsza, pokrewna wersja projektu tego samego autora. open-electronics.org/an-open-project-for-a-1000w-induction-heater

Topologia ZVS – oscylator Royera/Mazzilliego

- [Kaizer Power Electronics – Royer ZVS Induction Heater](https://kaizerpowerelectronics.dk/general-electronics/royer-induction-heater) – obszerny opis topologii, obliczenia obwodu LC, oscylogramy oraz cenne uwagi praktyczne dotyczące jej ograniczeń. kaizerpowerelectronics.dk/general-electronics/royer-induction-heater
- [RoboticWorx – High Voltage Generator & Induction Heater](https://roboticworx.io/blogs/projects/inducer-high-voltage-generator-induction-heater) – projekt ZVS z kompletem plików: Gerber, projekt KiCad i wykaz elementów (repozytorium GitHub podlinkowane w artykule). roboticworx.io/blogs/projects/inducer-high-voltage-generator-induction-heater
- [hardboiledfrog/smt-zvs-driver \(GitHub\)](https://github.com/hardboiledfrog/smt-zvs-driver) – wersja SMD oscylatora ZVS – projekt płytki, wariant zasilany z akumulatora ok. 700 W. github.com/hardboiledfrog/smt-zvs-driver
- [matebuteler/indheater \(GitHub\)](https://github.com/matebuteler/indheater) – nagrzewnica ZVS 1,4 kW oparta na projekcie Schematix – pliki płytki i wykaz materiałów. github.com/matebuteler/indheater
- [Instructables – Simple DIY Induction Heater With ZVS Driver](https://instructables.com/Simple-DIY-Induction-Heater-With-ZVS-Driver) – prosty przewodnik krok po kroku, dobry punkt wyjścia dla początkujących. instructables.com/Simple-DIY-Induction-Heater-With-ZVS-Driver

Mostek półkowy i konstrukcje z IGBT

- [Danyk.cz – Induction heating III \(mostek półkowy z IGBT\)](https://danyk.cz/induk3_en.html) – podwójny mostek półkowy z czterema IGBT STGW30NC60W sterowany układem IR2153 – schemat, zdjęcia, opis. danyk.cz/induk3_en.html
- [Danyk.cz – Induction heating I i II](https://danyk.cz/induken.html) – prostsze nagrzewnice z mostkiem i układem IR2153 zasilane wprost z sieci. danyk.cz/induken.html
- [Danyk.cz – IGBT halfbridge 20...200 kHz](https://danyk.cz/igbt3_en.html) – uniwersalny mostek półkowy IGBT z zabezpieczeniem nadprądowym – baza dla nagrzewnicy dużej mocy. danyk.cz/igbt3_en.html

Projekty i zestawy dostępne w polskich źródłach

Tematyka nagrzewnic indukcyjnych ZVS była wielokrotnie podejmowana w polskiej prasie elektronicznej. Poniżej zebrano projekty i zestawy, które stanowią bezpośrednie odpowiedniki konstrukcji omawianej w częściach A i B – mogą posłużyć czytelnikowi jako gotowe do zbudowania warianty albo materiał porównawczy.

- [Elektronika Praktyczna – „Domowa nagrzewnica indukcyjna” \(EP 10/2018\)](https://ep.com.pl/projekty/projekty-ep/12518-domowa-nagrzewnica-indukcyjna) – projekt nagrzewnicy w topologii ZVS (samowzbudny generator LC, IRFP260N), zasilanie 9...40 V DC, prąd do 40 A – najbliższy odpowiednik projektu z części A; dostępny w portalu EP.com.pl wraz z dokumentacją PDF. ep.com.pl/projekty/projekty-ep/12518-domowa-nagrzewnica-indukcyjna
- [AVT5645 – zestaw/płytki „Domowa nagrzewnica indukcyjna”](https://sklep.avt.pl/avt5645.html) – kit do samodzielnego montażu odpowiadający projektowi z EP 10/2018 (oscylator ZVS, 2× IRFP260N, 6× kondensator 330 nF) – dostępny w sklepie AVT w wersjach od samej płytki PCB po komplet elementów. sklep.avt.pl/avt5645.html
- [AVT2940 – zestaw „Nagrzewnica indukcyjna 1 kW” \(EdW 5/2010\)](https://sklep.avt.pl/manuals/AVT2940.pdf) – starszy projekt z „Elektroniki dla Wszystkich” w topologii mostka półkowego z układem IR2153 i tranzystorami IGBT – odpowiednik rozwiązania z punktu B.2; dokumentacja w archiwum serwisowym AVT. [serwis.avt.pl/manuals/AVT2940.pdf](https://sklep.avt.pl/manuals/AVT2940.pdf)
- [Elportal.pl – „Mininagrzewnica indukcyjna”](https://elportal.pl/projekty/pracownia-elektronika/6397-mininagrzewnica-indukcyjna) – miniaturowa nagrzewnica zasilana napięciem 24 V, oparta na zmodyfikowanej przetwornicy ZVS Mazzilliego (dwa dławiki zamiast odczepu uzwojenia) – projekt z „Elektroniki dla Wszystkich”. elportal.pl/projekty/pracownia-elektronika/6397-mininagrzewnica-indukcyjna
- [Elportal.pl – „Przetwornice indukcyjne cz. 28: jak działa przetwornica ZVS”](https://elportal.pl/projekty/pracownia-elektronika/6397-mininagrzewnica-indukcyjna) – odcinek kursu wyjaśniający zasadę działania oscylatora Royera/Mazzilliego – materiał teoretyczny uzupełniający część A. [elportal.pl – Przetwornice indukcyjne cz. 28](https://elportal.pl/projekty/pracownia-elektronika/6397-mininagrzewnica-indukcyjna)
- [Sklep AVT – moduł „Piec indukcyjny ZVS 12...48 V/20 A/1000 W”](https://sklep.avt.pl) – gotowy, zmontowany moduł nagrzewnicy ZVS wraz z cewką roboczą – odpowiednik konstrukcji z części A w postaci handlowej; przydatny do porównań i szybkiego startu. [sklep.avt.pl – Piec indukcyjny ZVS 12–48 V/20 A](https://sklep.avt.pl)

Noty aplikacyjne i materiały uzupełniające

- [ON Semiconductor \(Fairchild\) – AN-9012](https://onsemi.com/AN-9012) – „Induction Heating System Topology Review” – przegląd topologii przemysłowych, w tym quasi-rezonansowej; do pobrania w bazie not aplikacyjnych onsemi.com. [onsemi.com \(nota AN-9012\)](https://onsemi.com/AN-9012)
- [EEVblog – wątek „ZVS Induction Heater”](https://eevblog.com/forum/projects/zvs-induction-heater-2000w) – dyskusja konstruktorska o budowie i skalowaniu nagrzewnic ZVS dużej mocy. eevblog.com/forum/projects/zvs-induction-heater-2000w
- [Highvoltageforum.net](https://highvoltageforum.net) oraz [Elektroda.com](https://elektroda.com) – fora konstruktorskie z licznymi wątkami o nagrzewnicach ZVS i mostkach IGBT (m.in. budowy oparte na IRFP260N). highvoltageforum.net



Zegary hobbystyczne

Projekt przykładowy: zegar LED sterowany sygnałem GPS (AVT5522)

Zegar to jeden z tych projektów, które elektronik buduje przynajmniej raz – i do których wraca przez całe życie. Sposobów na odmierzanie i pokazywanie czasu jest dziesiątki: od liczników TTL bez procesora, przez konstrukcje z mikrokontrolerem i układem RTC, zegary synchronizowane sygnałem DCF77 lub GPS, aż po efektowne lampy Nixie i wyświetlacze analogowe. Omówienie tego tematu zawiera dwie części. W części A prezentujemy przykładowy projekt – zegar z dużym wyświetlaczem LED, który jako wzorzec czasu wykorzystuje odbiornik GPS. W części B poddajemy ten projekt komentarzowi redakcyjnemu, a następnie szerzej omawiamy rozwiązania spotykane w konstrukcjach zegarów i zestawiamy projekty pokrewne.

Część A. Projekt przykładowy Zegar ustawiany za pomocą GPS

Zadaniem zegara jest jak najdokładniejsze wskazywanie bieżącej godziny – najlepiej tak, aby to zegar informował nas o czasie, a nie my jego, korygując wskazania. Ale nawet zegar o bardzo dobrym „chodzie” wymaga okresowej korekty: ze względu na zmianę czasu letni/zimowy, a gdy ma wbudowany kalendarz – także na lata przestępne. Popularną metodą synchronizacji jest odbiornik DCF77; system jest sprawdzony, lecz podatny na zakłócenia, a duża antena ferrytowa przywodzi na myśl stary radioodbiornik. Opisywany zegar korzysta z rozwiązania nowocześniejszego – jako wzorzec czasu wykorzystuje system GPS, dzięki czemu można mu przypisać „kosmiczną dokładność”.

Co istotne, mimo przystosowania do współpracy z odbiornikiem GPS urządzenie może też pracować jako zwykły zegar ustawiany ręcznie – odbiornik nie jest elementem niezbędnym. Dzięki wyświetlaczowi LED o dużej wysokości cyfr jest to konstrukcja przeznaczona do montażu

Projekt: zegar z wyświetlaczem LED synchronizowany sygnałem GPS, oparty na mikrokontrolerze ATmega8 i układzie RTC MCP7940.

Źródło: artykuł „Zegar ustawiany za pomocą GPS”, Elektronika Praktyczna 9/2015. Zestaw dostępny w ofercie AVT jako AVT5522.

Dwa warianty zestawu: AVT5522/1 – z wyświetlaczem standardowym (cyfry 2 cm, sekundy 1,4 cm); AVT5522/2 – z wyświetlaczem „dużym” o wysokości cyfry 55 mm. Oba korzystają z tej samej płytki sterownika; różnią się płytką wyświetlacza i kilkoma elementami.

w miejscach publicznych: poczekalniach, halach, szkołach czy obiektach użyteczności publicznej.

Najważniejsze właściwości:

- wyświetlanie czasu (godziny, minuty, sekundy) oraz daty;
- opcjonalna synchronizacja czasu za pomocą odbiornika GPS (kupowanego oddzielnie);
- automatyczna zmiana czasu na letni i zimowy;
- możliwość wyświetlania temperatury (czujnik DS18B20);



Fotografia 1. Zegar z zielonym wyświetlaczem 55 mm (zestaw AVT5522/2/G)

- bateryjne podtrzymanie pamięci zegara RTC;
- wyświetlacz LED – wersja standardowa lub „duża” o wysokości cyfry 55 mm;
- zasilanie 12 V DC/0,2 A.

Zasada działania i synchronizacja czasu

Pracą urządzenia steruje mikrokontroler ATmega8 wraz z zawartym w jego pamięci programem. Pomiędzy punktami synchronizacji zegar odmierza czas za pomocą wyspecjalizowanego układu zegara czasu rzeczywistego (RTC) typu MCP7940. Układ ten pracuje w konfiguracji, w której na wyprowadzeniu MFP generowany jest przebieg prostokątny o częstotliwości dokładnie 1 Hz. Wyprowadzenie MFP połączono z wejściem przerwania INTO procesora – każde zbocze opadające powoduje wygenerowanie przerwania, a tym samym „tyknięcie” zegara. Ponieważ sygnał 1 Hz z układu RTC nie jest zsynchronizowany z czasem UTC, wskazania zegara mogą być spóźnione maksymalnie o jedną sekundę.

Synchronizacja z systemem GPS odbywa się po włączeniu zasilania, a następnie automatycznie co 3 godziny. Urządzenie oczekuje wówczas na właściwą ramkę **RMC** (ang. *recommended minimum GPS data*), przesyłaną przez większość odbiorników w standardzie NMEA-0183. Jeśli w ciągu 30 sekund poprawna ramka nie zostanie odebrana, sytuację tę traktuje się jako pracę bez synchronizacji – sygnalizuje ją krótkie miganie dwukropka na wyświetlaczu (ok. 1/3 s). Odebranie poprawnej ramki powoduje aktualizację wskazania oraz zawartości układu RTC i jest sygnalizowane długim miganiem dwukropka (ok. 2/3 s).

Dla zegara najistotniejszy jest czas UTC – wzorcowy czas uniwersalny ustalany na podstawie międzynarodowego czasu atomowego (TAI) i koordynowany względem ruchu obrotowego Ziemi. Jest to czas południka zerowego, dlatego aby zegar wskazywał czas lokalny, należy ustawić obowiązującą strefę czasową. Przykładowa ramka RMC, której parametry zestawiono w tabeli 1, ma postać:

\$GPRMC,220516,A,5133.82,N,00042.24,W,173.8,231.8,130694,004.2,W*70

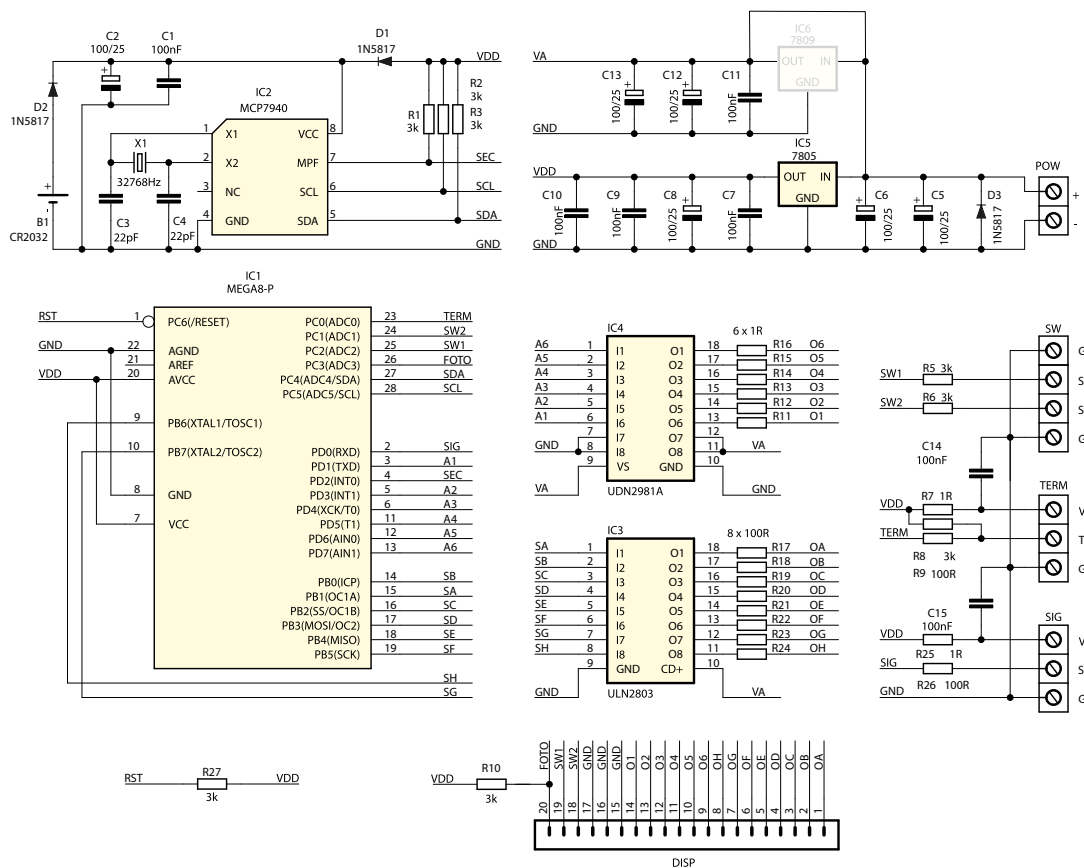
Tabela 1. Parametry zawarte w ramce RMC standardu NMEA-0183

Lp.	Parametr	Znaczenie
1	\$	znacznik początku ramki
2	GPRMC	typ ramki – RMC (recommended minimum GPS data)
3	220516	aktualna godzina w formacie HHMMSS (UTC)
4	A	aktualność danych: A – ważne, V – nieaktywne
5	5133.82	szerokość geograficzna
6	N	półkula: N – północna, S – południowa
7	00042.24	długość geograficzna
8	W	kierunek: W – zachodnia, E – wschodnia
9	173.8	prędkość obiektu w węzłach
10	231.8	kąt kursu w stopniach
11	130694	aktualna data w formacie DDMMYY
12	004.2	odchylenie magnetyczne
13	W	kierunek odchylenia magnetycznego
14	*70	suma kontrolna

Czas uniwersalny nie uwzględnia zmian na czas letni i zimowy. Informację o tym, który z nich obowiązuje, program uzyskuje, analizując datę zawartą w ramce RMC. Do obliczeń potrzebny jest jeszcze dzień tygodnia: w programie zapisano, że 1 stycznia 2000 roku przypadał w sobotę, więc znając bieżącą datę można wyliczyć liczbę dni, jakie upłynęły od tego punktu odniesienia, i wyznaczyć dzień tygodnia. Operowanie strukturą daty rozbity na rok, miesiąc, dzień, godzinę, minutę i sekundę przy dodawaniu i odejmowaniu godzin bardzo komplikuje program, dlatego przyjęto wygodniejszą reprezentację: czas względny – liczbę sekund, które upłynęły od 1 stycznia 2000 roku. Zmienna ma rozmiar 32 bitów, co w zupełności wystarcza do zliczania sekund do roku 2100. Po każdej sekundzie **czas względny** jest zwiększany o jeden, następnie korygowany o wartość wynikającą z ustawionej strefy czasowej, a jeśli obowiązuje czas letni – powiększany o 3600 sekund. Dopiero z tak przygotowanej wartości tworzona jest struktura daty i godziny prezentowana na wyświetlaczu.

Budowa układu

Zegar składa się z dwóch płytek: płytki głównej (sterownika) oraz płytki wyświetlacza. Schemat ideowy sterownika przedstawia rysunek 1.



Rysunek 1. Schemat ideowy płytki głównej zegara

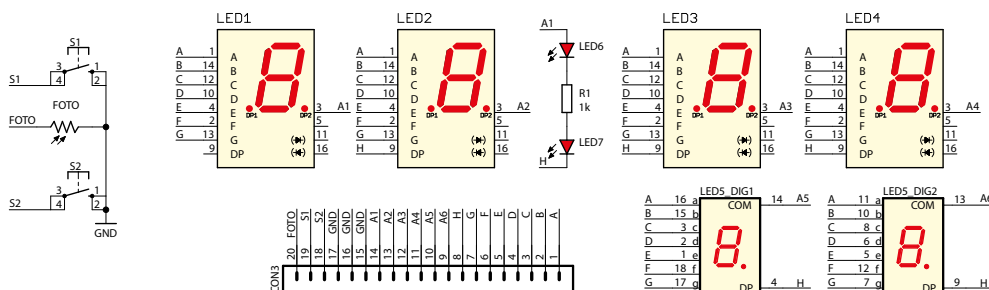
Sterownik

Centralnym elementem płytki głównej jest mikrokontroler ATmega8 (IC1). Odmierzanie czasu pomiędzy synchronizacjami powierzono układowi RTC MCP7940 (IC2), taktowanemu rezonatorem kwarcowym zegarkowym 32 768 Hz (X1); ciągłość pracy pamięci zegara przy zaniku zasilania zapewnia bateria litowa CR2032 (B1). Sygnał 1 Hz z wyprowadzenia MFP trafia na wejście INTO mikrokontrolera.

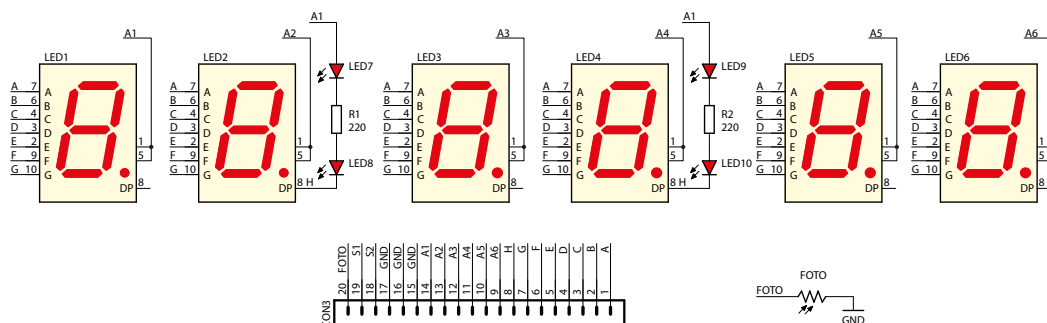
Wyświetlacze sterowane są zmodyfikowaną metodą multipleksową. Zasilanie od strony katod realizuje ośmiokrotny driver ujemnej szyny zasilania ULN2803 (IC3), natomiast anody zasila podobny driver, lecz sterujący biegunem dodatnim – UDN2981 (IC4; można zastosować odpowiednik, np. TD62783). Cyfry zapalają się kolejno; w danej chwili świeci zawsze tylko jedna, lecz proces przebiega na tyle szybko, że oko widzi cały

wyświetlacz. Modyfikacja polega na tym, że po każdym cyklu zaświecenia cyfry następuje krótszy cykl pełnego wygaszenia wszystkich wyświetlaczy – zapobiega to prześwitywaniu cyfr na sąsiednie pola, wynikającemu z czasu propagacji driverów. Zmieniając stosunek czasu świecenia do czasu wygaszenia, reguluje się zarazem intensywność świecenia. Parametrem sterującym jest sygnał z fotorezystora: im silniejsze oświetlenie otoczenia, tym jaśniej świeci wyświetlacz.

Zasilanie zapewniają dwa stabilizatory. Stabilizator 7805 (IC5) dostarcza napięcia +5 V dla mikrokontrolera i jego peryferiów. Drugi stabilizator (IC6, 7809) dostarcza +9 V dla wyświetlaczy w wersji standardowej. W wersji z dużym wyświetlaczem stabilizator IC6 zastępuje się zwrą (łącząc skrajne wyprowadzenia) i zasila urządzenie wyższym napięciem – maksymalnie +24 V. Prąd pojedynczego segmentu wyświetlacza nie może przekraczać 100 mA.



Rysunek 2. Schemat ideowy płytki wyświetlacza standardowego (zestaw AVT5522/1)



Rysunek 3. Schemat ideowy płytki wyświetlacza „dużego” (zestaw AVT5522/2)

Płytką wyświetlacza – dwa warianty

Ten sam sterownik obsługuje dwie różne płytki wyświetlacza, oferowane jako odrębne zestawy. W wariantcie standardowym (zestaw AVT5522/1) godziny i minuty pokazują wyświetlacze o wysokości cyfry 2 cm (LED1–LED4), a sekundy – mniejszy, dwucyfrowy moduł 1,4 cm; schemat ideowy tej płytki przedstawia rysunek 2.

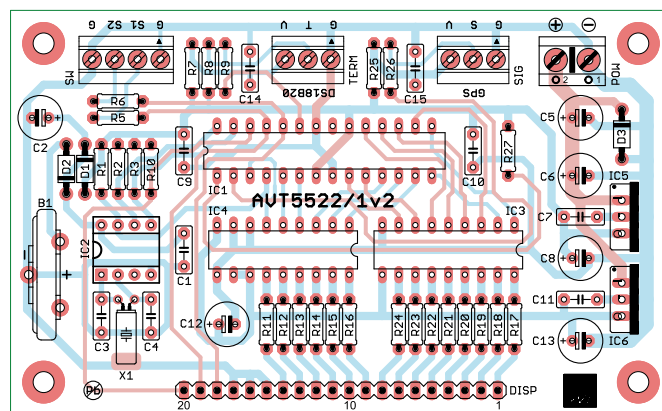
W wariantcie z wyświetlaczem „dużym” (zestaw AVT5522/2) zastosowano sześć identycznych wyświetlaczy o wysokości cyfry 5,6 cm (LED1–LED6); jego schemat ideowy przedstawia rysunek 3. W obu wariantach dwukropki rozdzielające pola tworzą diody LED (LED7–LED10), a płytkę wyświetlacza łączy się ze sterownikiem rzędem kątowych goldpinów. Wybór wariantu pociąga za sobą drobne różnice montażowe, opisane dalej – obecność stabilizatora IC6 oraz wartości rezystorów R15 i R16.

Montaż i uruchomienie

Płytki zaprojektowano dla elementów przewlekanych, więc montaż nie powinien sprawić trudności nawet osobom z mniejszym doświadczeniem. Elementy wlotowuje się w kolejności gabarytowej – od najniższych do najwyższych. Złącze CON1 umożliwia programowanie mikrokontrolera w układzie, lecz nie musi być montowane, ponieważ procesor

w zestawie jest dostarczany zaprogramowany. Schemat montażowy płytki głównej przedstawia rysunek 4.

Na płytce wyświetlacza w pierwszej kolejności montuje się listwę szpilek goldpin – po stronie lutowania – dzięki czemu obie płytki można następnie łatwo połączyć. Same wyświetlacze warto zamontować kilka milimetrów ponad powierzchnią płytki: pozwala to wsunąć nakrętki za otwory montażowe i solidnie skrócić zestaw z użyciem śrub oraz tulejek dystansowych. Po zmontowaniu z zaprogramowanym mikrokontrolerem zegar jest gotowy do pracy.

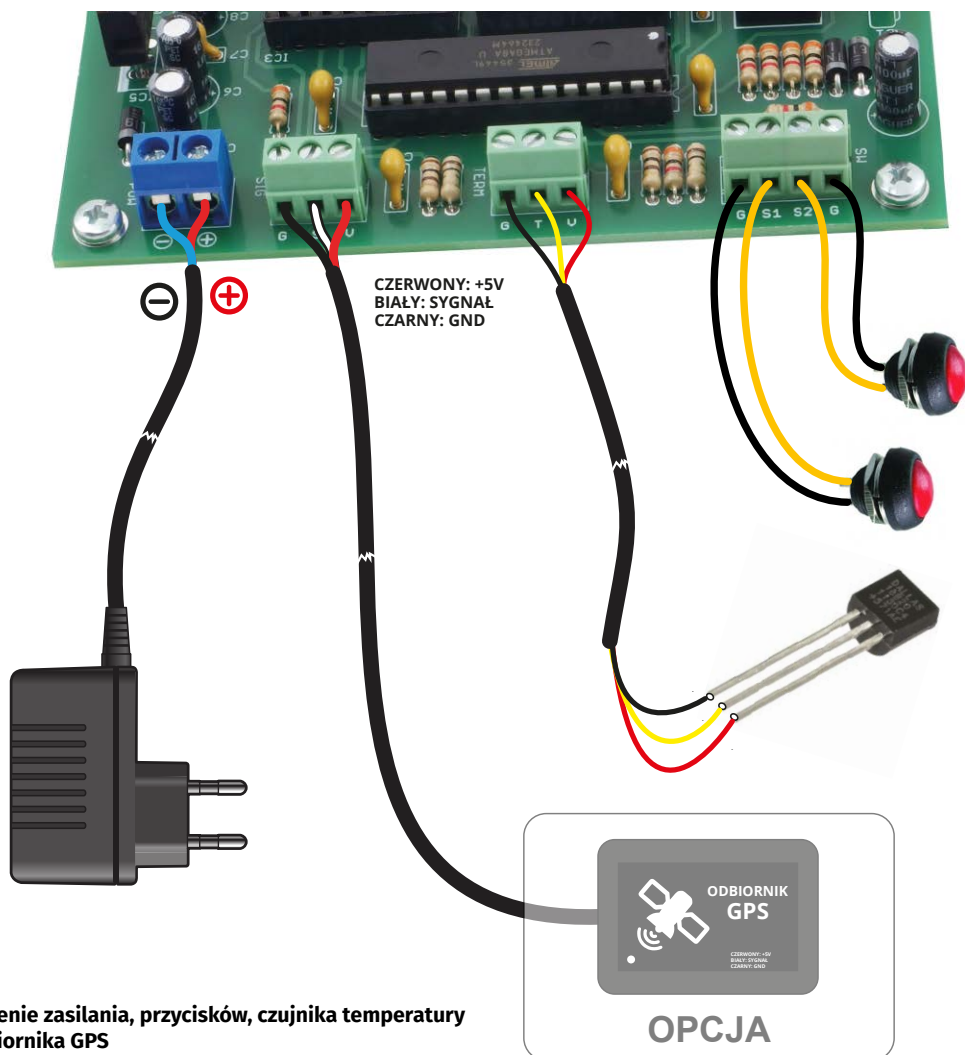


Rysunek 4. Schemat montażowy płytki głównej zegara. Zwróć uwagę na miejsce stabilizatora IC6 w wersji z dużym wyświetlaczem

Tabela 2. Wykaz elementów zegara. Czujnik temperatury DS18B20 i odbiornik GPS nie wchodzi w skład każdej wersji zestawu

Płytką główną	
Rezystory	R1...R3, R5, R6, R8, R10, R27: 3 kΩ; R9, R15...R24, R26: 100 Ω; R7, R11...R16, R25: 1 Ω
Kondensatory	C1, C7, C9...C11, C14, C15: 100 nF; C2, C5, C6, C8, C12, C13: 100 μF/25 V; C3, C4: 22 pF
Półprzewodniki	D1...D3: 1N5817 (lub podobne); IC1: ATmega8 (zaprogramowany); IC2: MCP7940; IC3: ULN2803; IC4: UDN2981 lub np. TD62783; IC5: 7805; IC6: 7809 lub zwora; TERM: DS18B20
Pozostałe	X1: kwarc zegarkowy 32 768 Hz; B1: bateria litowa 3 V do druku; CON1: nie montować; DISP: goldpin kątowy 1×20; POW: ARK2/500; SW, TERM, SIG: ARK3/500
Płytką wyświetlacza standardowego	
Elementy	R1: 1 kΩ; S1, S2: przyciski; FOTO: fotorezystor GL5616 (1M); LED1...LED4: LED-AS8012; LED5, LED6: LED-AD5624; LED7...LED10: dioda LED 3 mm; DISP: goldpin kątowy 1×20
Płytką wyświetlacza „dużego” (opcjonalnego)	
Elementy	R1, R2: 220 Ω; FOTO: fotorezystor GL5616 (1M); LED1...LED6: LED-AS23011; LED7...LED10: dioda LED 5 mm; DISP: gniazdo goldpin 1×20

Uwaga redakcyjna: rezystory R15 i R16 mają różną wartość zależnie od wariantu – 100 Ω dla wyświetlacza standardowego (wyrównanie jasności mniejszego pola sekund) i 1 Ω dla wyświetlacza dużego. W wersji standardowej montuje się stabilizator IC6, w wersji z dużym wyświetlaczem zastępuje się go zworą



Rysunek 5. Podłączenie zasilania, przycisków, czujnika temperatury i opcjonalnego odbiornika GPS

Podłączenie zegara

Zasilanie 12 V DC doprowadza się do złącza POW. Czujnik temperatury DS18B20 dołącza się do złącza TERM, a przyciski S1 i S2 – do złącza SW (wygodnie wyprowadzić je przewodami, np. na tylną ściankę obudowy). Odbiornik GPS dołącza się do złącza SIG: przewód masy do zacisku G, sygnału do zacisku S i zasilania do zacisku V. Sposób połączenia całości pokazuje rysunek 5.

Odbiornik GPS musi mieć interfejs UART o poziomie napięcia 3,3 V lub 5 V, z parametrami transmisji: 8 bitów danych, brak parzystości, jeden bit stopu. Obsługiwane prędkości to 4800, 9600, 14400 i 19200 b/s. Wymagania spełnia

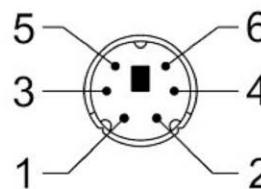
m.in. odbiornik MARS600 zbudowany na module u-blox 6 (rysunek 6) – ma zintegrowaną antenę i zamknięty jest w niewielkiej obudowie z magnesem. Odbiornik należy umieścić tak, aby „widział” niebo. Warto pamiętać, że odbiornik GPS nie jest niezbędny do pracy zegara; przy pracy bez niego zegar wymaga jednorazowego ustawienia czasu i okresowej kontroli wskazań.

Obsługa

Do obsługi zegara służą dwa przyciski – S1 i S2. Krótkie naciśnięcie S1 wyświetla datę w formacie DD:MM:RR (po ok. 5 sekundach zegar wraca do wskazywania czasu).



Pin	Signal
1	GND
2	VCC 5.0V
3	N.C.
4	TTL RX
5	N.C.
6	TTL TX



Mars600-mini-T

Rysunek 6. Przykładowy odbiornik GPS MARS600 (moduł u-blox 6) wraz z opisem wyprowadzeń złącza

Krótkie naciśnięcie S2 pokazuje temperaturę odczytaną z czujnika. Dłuższe (ok. 3 s) przytrzymanie S1 uruchamia tryb ręcznego ustawiania daty i czasu: edytowane pole miga na przemian z kursorem, przycisk S2 zwiększa wartość, S1 – zmniejsza, a dłuższe przytrzymanie przesuwa kursor na kolejne pole. Dłuższe przytrzymanie S2 wymusza synchronizację na żądanie – zegar zaczyna śledzić sygnał z odbiornika GPS i aktualizuje czas po odebraniu ramki RMC.

Część parametrów ustawia się przy włączaniu zasilania. Włączenie z wciśniętym S1 otwiera ustawienia strefy czasowej – w zakresie od -14 do +14 godzin, ze skokiem 1 godziny (program nie obsługuje stref typu +3:30). Wartość domyślna to +1, więc dla Polski nie trzeba jej zmieniać; osobne pole włącza lub wyłącza automatyczną zmianę czasu letni/zimowy. Włączenie zasilania z wciśniętym S2 otwiera ustawienia dodatkowe: prędkość transmisji odbiornika GPS (0 – 4800, 1 – 9600, 2 – 14 400, 3 – 19 200 b/s; domyślnie 1) oraz status automatycznego, cyklicznego wyświetlania temperatury.

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 zegarach amatorskich – od logiki bez procesora po synchronizację sieciową. Na koniec zestawiamy projekty pokrewne i literaturę.

Komentarz redakcyjny do projektu AVT5522

Projekt obroni się nawet po dekadzie. Mocną stroną jest podział na dwie płytki i wynikająca z niego skalowalność – ten sam sterownik obsługuje wyświetlacz standardowy i „duży”, z napięciem anodowym doborczym do potrzeb (do 24 V). Automatyczna obsługa strefy czasowej i przejść letni/zimowy, bateryjne podtrzymanie układu RTC oraz montaż wyłącznie z elementów przewlekanych czynią go konstrukcją użytkową i przyjazną w budowie. Poniżej zbieramy uwagi, które warto mieć na uwadze przy ewentualnej modernizacji.

Błąd rzędu jednej sekundy – autorzy projektu rzetelnie odnotowują, że przebieg 1 Hz z wyprowadzenia MFP układu RTC nie jest zsynchronizowany w fazie z sekundą UTC – między punktami synchronizacji wskazanie sekund może być przesunięte maksymalnie o 1 s. To świadomy kompromis na rzecz prostoty. Ścieżką rozwoju jest wykorzystanie wyjścia 1PPS odbiornika GPS (dostępnego m.in. w modułach u-blox) do zdyscyplinowania lub bezpośredniego taktowania sekundnika; AVT5522

korzysta wyłącznie z ramki NMEA przesyłanej przez UART, pomijając 1PPS.

Reguły czasu letniego zaszyte w firmware – logika przejść (ostatnia niedziela marca i października) jest zapisana na stałe w programie. Dla wdrożeń poza strefą obowiązywania tych reguł – lub gdyby reguły uległy zmianie – konieczna jest modyfikacja oprogramowania. Rozwiązanie tablicowe lub konfigurowalne byłoby odporniejsze. Drobnym, lecz wartym wzmianki ograniczeniem jest też 32-bitowy licznik sekund liczonych od 2000 roku, wymagający korekty po roku 2100.

Sterowanie multipleksowe a driver prądowy – wyświetlacze pracują w multipleksie, z driverami ULN2803 i UDN2981 oraz rezystorami szeregowymi; jasność reguluje się stosunkiem czasu świecenia do wygaszenia. Przy sześciu dużych cyfrach współczynnik wypełnienia na pole jest niski, więc prąd szczytowy segmentu jest wysoki – stąd limit 100 mA. Dedykowany driver o stałym prądzie (np. z rodziny sterowników LED z wyjściami prądowymi) zapewniłby stabilną jasność niezależną od taktowania odświeżania i odciążał mikrokontroler.

Odbiornik GPS – uwagi praktyczne – zegar wymaga odbiornika z interfejsem UART 3,3 V lub 5 V; pokazany w projekcie moduł MARS600 oparty jest na układzie u-blox 6, dziś już archaicznym. Współczesne, tańsze moduły (typu NEO-6M/7M/8M i nowsze) spełniają te same wymagania i są w pełni zamienniki. Należy jedynie zachować zgodność poziomów napięć i parametrów transmisji.

Naturalna alternatywa: ESP32 i synchronizacja NTP – gdyby projektować ten zegar od nowa, w wielu zastosowaniach zamiast GPS sięgnięto by po mikrokontroler z Wi-Fi (ESP32, ESP8266) i synchronizację protokołem NTP/SNTP. Eliminuje to odbiornik, antenę i wymóg „widoku nieba”, a strefę czasową i reguły DST można pobierać z bazy stref. GPS pozostaje jednak przewagą tam, gdzie brak sieci – w instalacjach odległych czy mobilnych.

Jak odmierza się i synchronizuje czas – przegląd rozwiązań Zegary cyfrowe na układach logicznych (bez mikrokontrolera)

Zanim mikrokontroler stał się elementem oczywistym, zegar budowano wyłącznie na licznikach i bramkach. Schemat blokowy jest do dziś pouczający: generator wzorcowy 1 Hz (rezonator kwarcowy 32 768 Hz z łańcuchem dzielników, np. CD4060, albo dzielona częstotliwość sieci 50 Hz), kaskada liczników dzielących sekundy przez 60, minuty przez 60 i godziny przez 12 lub 24, oraz dekodery BCD → 7 segmentów (74xx47, CD4511). Liczniki dekadowe (7490, CD4017, CD4518) z logiką resetu realizują podział, który nie jest potęgą dwójki. Konstrukcja taka jest w pełni deterministyczna, nie wymaga firmware'u i – co cenne

dydaktycznie – pokazuje wprost, jak z jednej częstotliwości wzorcowej powstaje wskazanie czasu.

Mikrokontroler, układ RTC i biblioteki

Projekt przykładowy reprezentuje podejście, które zdominowało konstrukcje amatorskie: mikrokontroler odpowiada za logikę i wyświetlanie, a odmierzanie czasu powierza się wyspecjalizowanemu układowi RTC. Rodzina takich układów jest szeroka – od prostego DS1302, przez popularny DS1307, po DS3231 z wbudowanym, kompensowanym termicznie rezonatorem (TCXO) o dokładności rzędu ± 2 ppm. Zastosowany w AVT5522 układ MCP7940 udostępnia dodatkowo programowalne wyjście częstotliwościowe, wykorzystane tu jako źródło przerwania 1 Hz. Układ RTC zwykle samodzielnie obsługuje kalendarz wraz z latami przestępnymi i dniem tygodnia, a podtrzymanie bateryjne utrzymuje czas po zaniku zasilania. Po stronie oprogramowania dojrzałe biblioteki (np. RTCLib) sprowadzają obsługę do kilku wywołań. Warto odnotować, że najnowsze mikrokontrolery integrują układ czasu rzeczywistego na strukturze – przykładowo płytki Arduino UNO R4 udostępniają wewnętrzny RTC z obsługą alarmów, co eliminuje osobny element.

Synchronizacja: DCF77, GPS, 1PPS i NTP

Lokalny oscylator, choćby dobry, dryfuje – stąd potrzeba okresowej synchronizacji z wzorcem zewnętrznym. Historycznie pierwszym powszechnie dostępnym rozwiązaniem był DCF77: nadawany z Mainflingen sygnał 77,5 kHz, modulowany amplitudowo, przenoszący jeden bit na sekundę i pełną ramkę czasu w ciągu minuty. System jest dojrzały i darmowy, lecz podatny na zakłócenia i wymaga sporej anteny ferrytowej. GPS, zastosowany w projekcie przykładowym, czerpie czas z zegarów atomowych satelitów; sama ramka NMEA daje dokładność rzędu sekundy, natomiast osobne wyjście 1PPS pozwala – po odfiltrowaniu jittera – zdyscyplinować lokalny oscylator z dokładnością mikrosekundową. Najnowszym, a dziś często najwygodniejszym wariantem jest synchronizacja sieciowa:

mikrokontroler z Wi-Fi pobiera czas protokołem NTP/SNTP z hierarchii serwerów. Metody zestawia tabela 3.

Wyświetlacze specjalne: Nixie, VFD i wskazania analogowe

Sposób prezentacji czasu bywa dla konstruktora ważniejszy niż sama elektronika sterująca. Obok klasycznych wyświetlaczy 7-segmentowych LED – jak w projekcie przykładowym – popularne są lampy Nixie. Te elementy z zimną katodą wymagają napięcia rzędu 170...200 V, więc zegar Nixie zawiera przetwornicę podwyższającą i wysokonapięciowe sterowniki (historyczne układy typu 74141/K155ID1 lub matryce tranzystorowe), a projekt musi traktować poważnie bezpieczeństwo pracy z wysokim napięciem. Pokrewną grupą są wyświetlacze próżniowe VFD. Odrębny kierunek to reinterpretacja wskazania analogowego: pierścień 60 diod LED odwzorowujący tarczę albo – w wersji skrajnej – analogowe mierniki wskazówkowe użyte jako wskaźniki godzin i minut. Każde z tych rozwiązań to osobny problem konstrukcyjny: sterowanie HV, multipleksowanie dużej liczby źródeł światła bądź kalibracja miernika.

Zegar jako projekt: kity i aspekt warsztatowy

Dla wielu czytelników zegar jest przede wszystkim pretekstem – do nauki lutowania, czytania schematu, projektowania obudowy. Oferta zestawów rozciąga się od prostych konstrukcji edukacyjnych po kity rozbudowane (podświetlenie RGB, aktualizacja oprogramowania przez USB, sygnalizacja dźwiękowa). Projekt przykładowy, dostępny jako zestaw AVT w wariantach od samej płytki po wersję zmontowaną, mieści się w nurcie dojrzałych konstrukcji użytkowych. Niezależnie od stopnia złożoności pozostają stałe wyzwania warsztatowe: ergonomia (rozmieszczenie przycisków bez psucia estetyki frontu – w AVT5522 rozwiązane wyprowadzeniem przycisków przewodami), czytelność (kontrast wyświetlacza, filtr optyczny) oraz pewne mocowanie mechaniczne płytek.

Tabela 3. Porównanie metod odmierzania i synchronizacji czasu stosowanych w zegarach amatorskich

Metoda	Wzorec czasu	Dokładność	Zalety	Ograniczenia
Oscylator lokalny (RTC)	wewnętrzny rezonator kwarcowy	$\pm$ kilka s/dobę; ± 2 ppm dla TCXO (DS3231)	prostota, brak zależności od sygnału zewnętrznego	dryf, konieczna ręczna korekta
DCF77	zegar atomowy PTB (Mainflingen)	rzędu 1 s	dojrzały, darmowy, dostępny w Europie Środkowej	podatny na zakłócenia, duża antena, ograniczony zasięg
GPS (ramka NMEA)	zegary atomowe satelitów	rzędu 1 s	„kosmiczna” dokładność, zasięg globalny	wymaga widoku nieba, opóźnienie ramki na UART
GPS z sygnałem 1PPS	jw., z impulsem sekundowym	poniżej 1 μ s	bardzo wysoka, dyscyplinuje oscylator lokalny	złożony firmware, obsługa wejścia 1PPS
NTP/SNTP (Wi-Fi)	hierarchia serwerów (stratum)	rzędu milisekund	brak dodatkowego sprzętu RF, strefy z bazy danych	wymaga sieci Wi-Fi i dostępu do Internetu

Od logiki TTL po synchronizację sieciową – zegar pozostaje projektem, do którego konstruktor wraca na każdym etapie swojej drogi, za każdym razem znajdując frajdę w odkrywaniu nowych problemów inżynierskich.

WM

Projekty pokrewne i literatura

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

Zegary bez mikrokontrolera

- [Electronics-course.com](https://electronics-course.com) – 12H/24H Digital Clock Circuit. Omówienie bloku generatora 1 Hz oraz tańcucha dzielników sekund, minut i godzin (na licznikach 74xx93) – dobra podbudowa teoretyczna. <https://electronics-course.com/digital-clock>
- [PCBWay](https://www.pcbway.com/project/shareproject/24_Hour_Digital_Clock_77a8b997.html) – 24 Hour Digital Clock (Share Project). Zegar 24-godzinny bez mikrokontrolera, na układach CMOS, z kompletnym schematem i płytką – współczesny odpowiednik rozwiązań TTL/CMOS. https://www.pcbway.com/project/shareproject/24_Hour_Digital_Clock_77a8b997.html

Mikrokontroler, RTC i zegary programowalne

- [Instructables](https://www.instructables.com/Make-A-Digital-Clock-From-Scratch/) – Make a Digital Clock From Scratch. Zegar na mikrokontrolerze z wyświetlaczami 7-segmentowymi – opis połączeń i interfejsu ustawiania czasu. <https://www.instructables.com/Make-A-Digital-Clock-From-Scratch/>
- [Instructables](https://www.instructables.com/DIY-Digital-Clock-Using-ATmega-328p-Arduino-IC-RTC/) – DIY Digital Clock Using ATmega328p, RTC DS3231 and Seven Segment Displays. Projekt modułowy: mikrokontroler z układem DS3231 i wyświetlaczami 7-segmentowymi – z alarmem, datą i pomiarem temperatury. <https://www.instructables.com/DIY-Digital-Clock-Using-ATmega-328p-Arduino-IC-RTC/>
- [Arduino](https://docs.arduino.cc/tutorials/uno-r4-minima/rtc/) – UNO R4 Minima – Real-Time Clock (dokumentacja oficjalna). Konfiguracja wbudowanego RTC, ustawianie czasu startowego, alarmy i odczyt – przykład integracji RTC na strukturze mikrokontrolera. <https://docs.arduino.cc/tutorials/uno-r4-minima/rtc/>
- [EXP-Tech](https://exp-tech.de/en/blogs/blog/arduino-tutorial-real-time-clock-rtc) – Arduino Tutorial: Real Time Clock (RTC). Praktyczne omówienie układu DS3231 na magistrali I²C – podłączenie, biblioteka, odczyt znaczników czasu. <https://exp-tech.de/en/blogs/blog/arduino-tutorial-real-time-clock-rtc>

Synchronizacja czasu

- [Chris Dzombak](https://www.dzombak.com/blog/2021/12/building-the-atomic-clock-i-ve-always-wanted/) – Building the atomic clock I've always wanted. Zegar zbudowany wokół zewnętrznego wzorca atomowego – praktyczne problemy integracji impulsu sekundowego z mikrokontrolerem. <https://www.dzombak.com/blog/2021/12/building-the-atomic-clock-i-ve-always-wanted/>

Wyświetlacze specjalne

- [SparkFun](https://news.sparkfun.com/2764) – DIY Nixie Tube Clock Build. Budowa zegara na lampach Nixie od strony praktycznej – pozyskanie lamp, montaż i wykonanie obudowy. <https://news.sparkfun.com/2764>
- [Instructables](https://www.instructables.com/SMD-Nixie-Clock/) – SMD Nixie Clock. Współczesna interpretacja zegara Nixie w technologii SMD – projekt PCB, przetwornica HV, montaż i firmware. <https://www.instructables.com/SMD-Nixie-Clock/>
- [All About Circuits](https://www.allaboutcircuits.com/projects/build-your-own-clock-with-analog-dials-part-1/) – Build Your Own Clock With Analog Dials (cz. 1–3). Zegar wykorzystujący trzy analogowe mierniki wskazówkowe jako wskaźniki czasu – przykład nietypowego interfejsu prezentacji. <https://www.allaboutcircuits.com/projects/build-your-own-clock-with-analog-dials-part-1/>

Baza projektów

- [Electronics For You](https://www.electronicshobby.com/category/electronics-projects/hardware-diy) – Hardware DIY – baza projektów. Stale aktualizowany zbiór projektów konstrukcyjnych, w tym rozmaitych zegarów – ogólne źródło inspiracji warsztatowej. <https://www.electronicshobby.com/category/electronics-projects/hardware-diy>

Zegary w ofercie kitów AVT

Opisany w części A zegar jest dostępny w sklepie AVT jako gotowy zestaw do samodzielnego montażu.

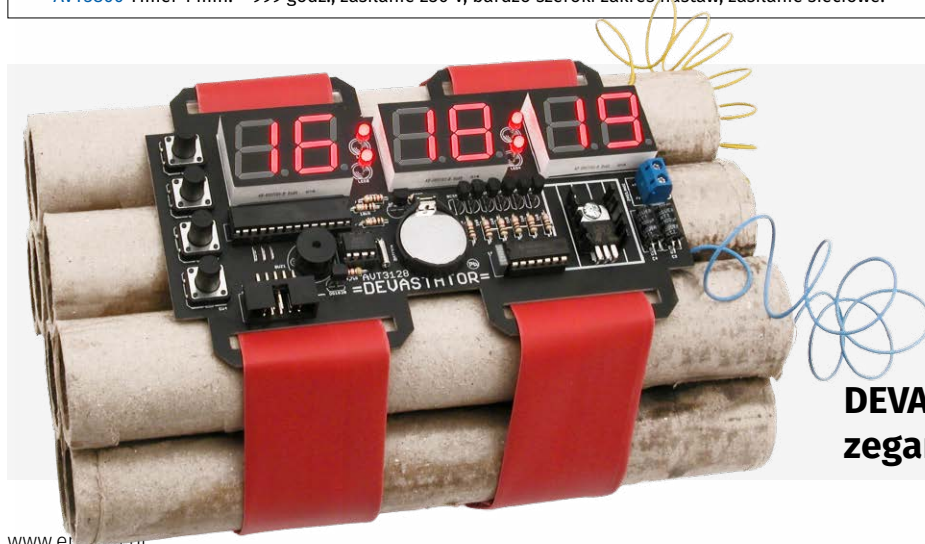
Poniżej zestawiamy całą aktualną ofertę zegarów oraz timerów ze sklepu AVT (sklep.avt.pl)

Zegary

- [AVT5522/1](#) Zegar ustawiany za pomocą GPS, wyświetlacz czerwony 20 mm. Projekt przykładowy z części A – wariant standardowy.
- [AVT5522/2](#) Zegar ustawiany za pomocą GPS, wyświetlacz 55 mm. Projekt przykładowy z części A – wariant „duży”; wyświetlacz czerwony lub zielony.
- [AVT5377](#) Stoper dla 5 zawodników – zegar i licznik zdarzeń, 2 wyświetlacze. rozbudowane odmierzenie czasu i zliczanie zdarzeń.
- [AVT3128](#) DEVASTATOR – „bombowy” zegarek. Efektowny zegar w konwencji gadżetowej.
- [AVT3132](#) Prosty zegar LED. Dostępny z wyświetlaczem czerwonym, zielonym, żółtym i białym;
- [AVT5805](#) Zegar RTC z kalendarzem, budzikiem i stoperem. Wyświetlacz OLED, rozbudowane funkcje czasowe.

Timery i wyłączniki czasowe

- [AVT1821](#) Timer ON/OFF. Prosty cykliczny włącznik czasowy.
- [AVT1995](#) Precyzyjny timer 1 s – 99 min. Nastawa od pojedynczych sekund.
- [AVT3200](#) Uniwersalny timer 0 – 99 min. Do typowych zastosowań warsztatowych i domowych.
- [AVT5638](#) Miniaturowy przypominacz. Kompaktowy układ sygnalizacji upływu czasu.
- [AVT5707](#) Wyłącznik czasowy 1 – 999 min. Szeroki zakres nastaw czasu wyłączenia.
- [AVT5899](#) Timer do pompy obiegowej. Sterowanie czasowe pracy pompy.
- [AVT3143](#) Nakręcany minutnik, mechaniczny minutnik kuchenny.
- [AVT5800](#) Timer 1 min. – 999 godz., zasilanie 230 V, bardzo szeroki zakres nastaw, zasilanie sieciowe.



<https://tiny.pl/8-2ycshx8>

DEVASTATOR, czyli bombowy zegarek (AVT3128)



W ofercie AVT*

AVT6098

Najważniejsze parametry:

- odliczanie czasu zadanego w przedziale 1...99 minut z rozdzielczością 1 minuty
- reprezentacja upływającego czasu w postaci 64 punktów świetlnych „przesypujących się” z góry na dół
- możliwość wstrzymania odliczania oraz jego wznowienia, jak również szybkie wyzerowanie już trwającego odliczania
- sygnalizacja dźwiękowa rozpoczęcia odliczania, jego wstrzymania oraz zakończenia
- zapamiętywanie ustawionego czasu na wypadek zaniku zasilania
- wyświetlacz w postaci dwóch dużych matryc LED 8×8
- pobór prądu około 120 mA
- zasilanie napięciem stałym 5 V poprzez gniazdo USB typu B lub listwę zaciskową

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

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

Projekty pokrewne na stronie www.ep.com.pl

(aktywne linki do artykułów):

- Magiczna kostka LED RGB na bazie RP2040
- Żarówka LED
- Regulowany, akumulatorowy zasilacz LED o mocy 3 W
- Regulator jasności diod power LED
- Iluminofonia LED RGB
- Gra „Kto pierwszy ten lepszy”
- Gra PONG na bazie Arduino i wyświetlacza LED
- Gra elektroniczna Sudoku

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

Elektroniczna klepsydra LED

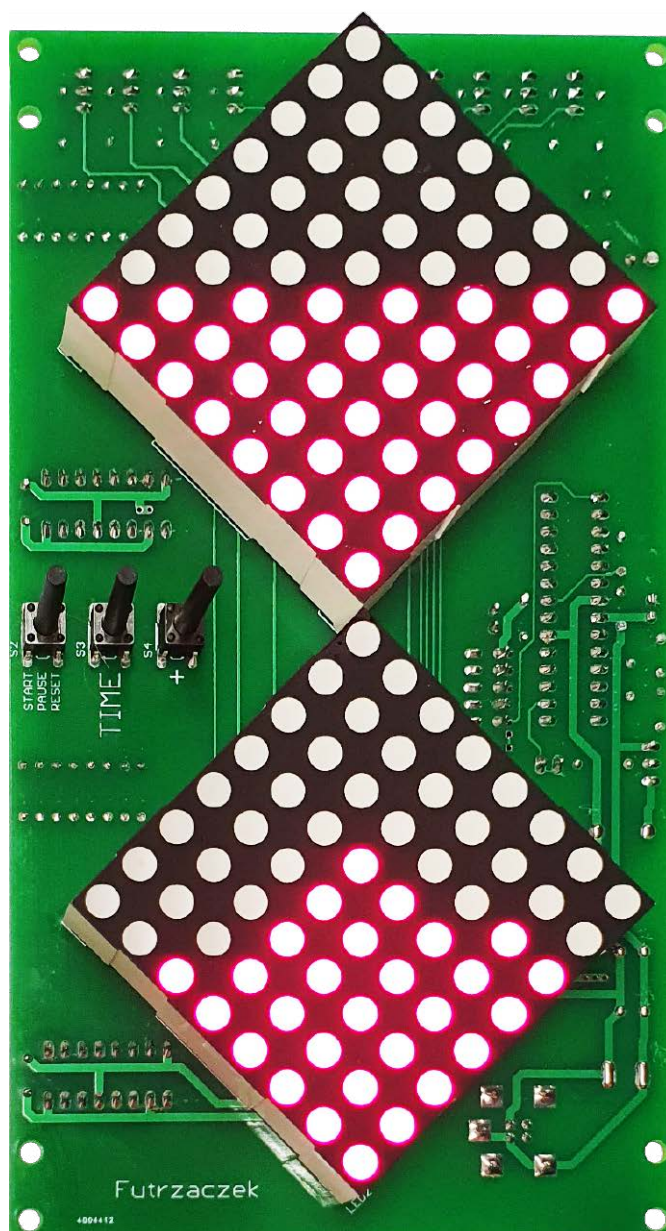
Czy przyrząd służący do odmierzania czasu może być niecodzienny, futurystyczny, przyciągający uwagę? Może. A czy może nawiązywać jednocześnie do historycznych rozwiązań w tej dziedzinie? Również może. Oto klepsydra, z której piasek nigdy się nie wysypie!

Minutnik w coraz rzadziej spotykanej formie – klepsydry. W porównaniu z jego tradycyjnym, analogowym pierwowzorem, ma trzy istotne funkcjonalności: można łatwo zmieniać czas do odliczenia, sygnalizuje dźwiękiem zakończenie odliczenia oraz odmierzanie czasu można w każdej chwili wstrzymać oraz wznowić, jak również wyzerować w okamgnieniu. Zwykle klepsydry tego nie potrafią! Zamiast ziarenek piasku i szklanej obudowy w układzie mamy kawałek elektroniki osadzony na płytce drukowanej. Imitacja pojemniczków z piaskiem są dwie kwadratowe matryce LED, ułożone jedna nad drugą i obrócone o 45° względem krawędzi płytki. Na nich jest wyświetlany zadany czas (w postaci liczby dwucyfrowej, po jednej cyfrze na matrycy) oraz wspomniane już kropki, kiedy układ wejdzie we właściwą fazę odmierzania czasu.

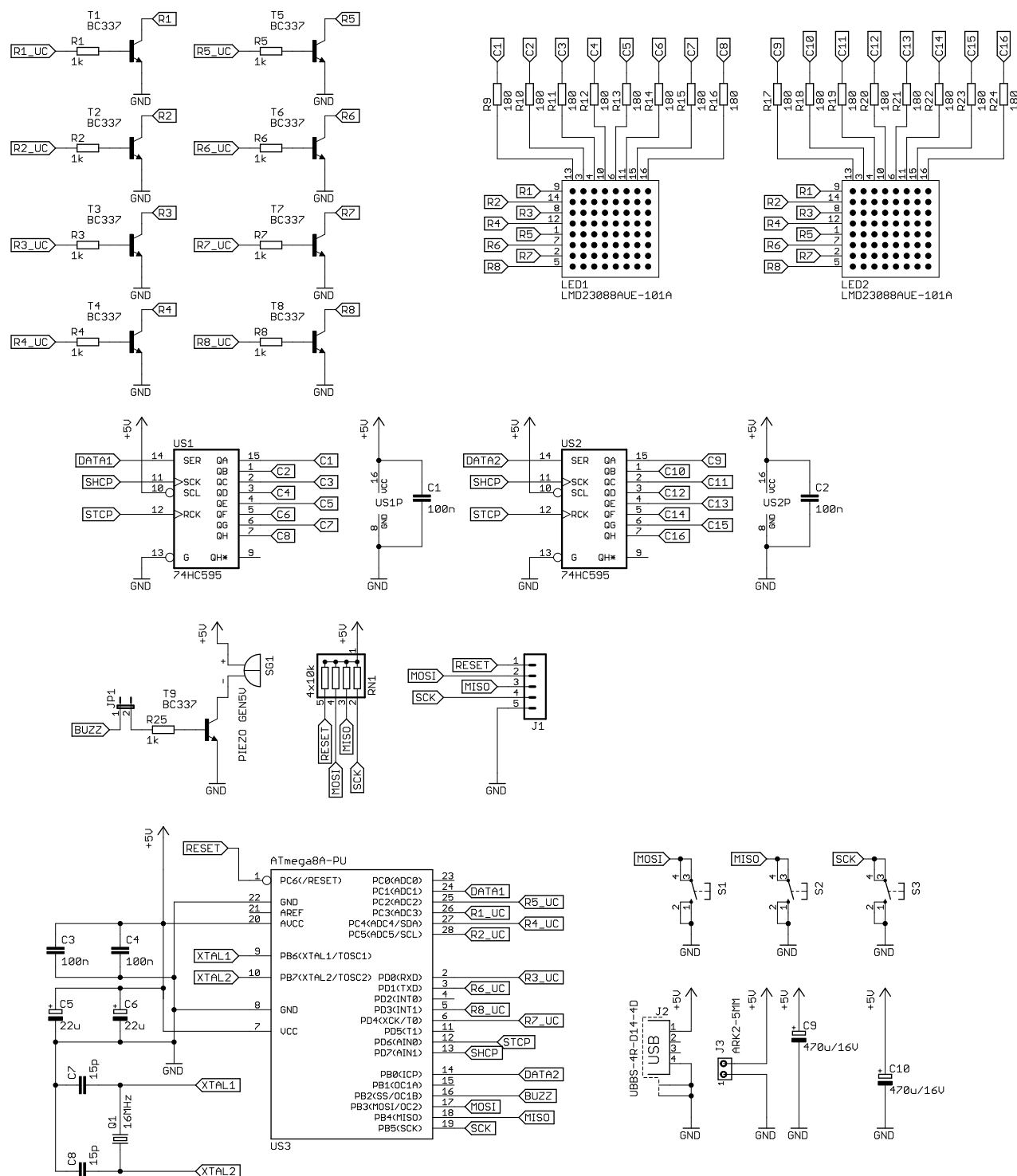
Taka klepsydra nigdy nie spadnie ze stołu, trudniej jest również oszukiwać w przeróżnych grach i zabawach – zwykłą klepsydre można ustawić ukośnie, by wolniej przesypywał się w niej piasek. Dźwiękowa sygnalizacja zakończenia odliczania utrudnia przegapienie końca swojej tury, a jeżeli ustalenia gry ulegną zmianie i czas dla każdego uczestnika będzie musiał zostać zmieniony, nie niesie to najmniejszych trudności. Wszak elektronika potrafi wszystko. No dobra, prawie wszystko.

Budowa

Schemat ideowy omawianego układu znajduje się na rysunku 1. Głównym podzespołem, zawiadującym jego pracą



jest mikrokontroler typu ATmega8A-PU z 8-bitowym rdzeniem AVR. Jego rdzeń jest taktowany sygnałem o częstotliwości 16 MHz, dla którego wzorcem jest rezonator kwarcowy



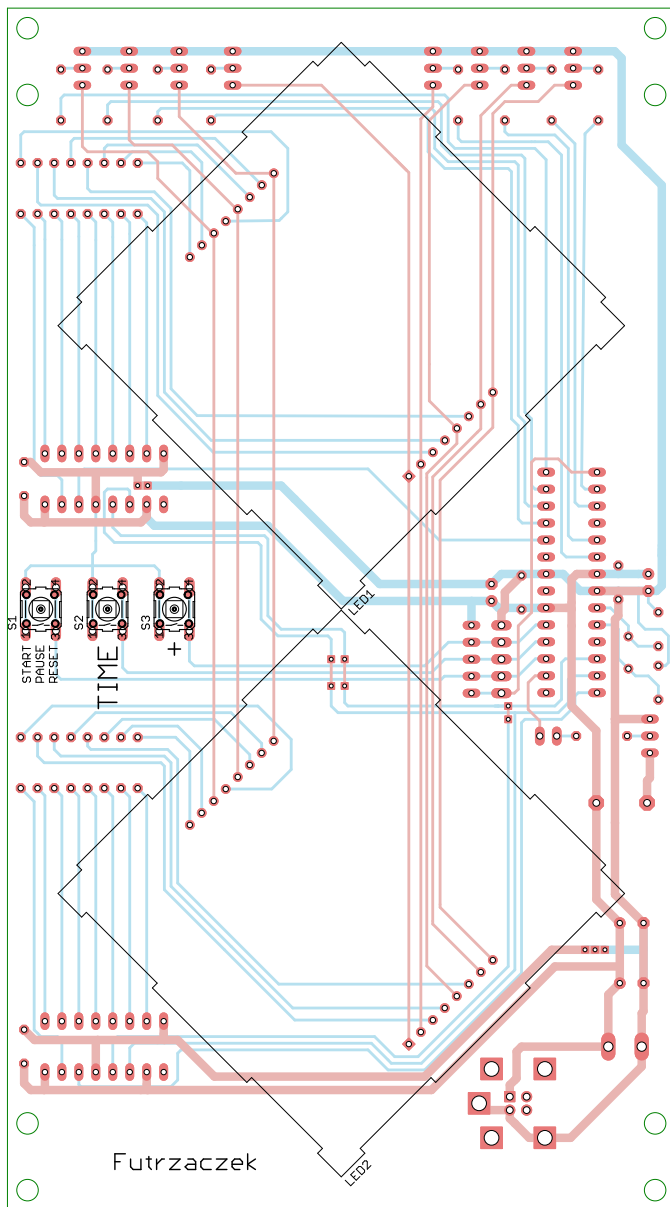
Rysunek 1. Schemat ideowy elektronicznej klepsydry LED

Q1. Wbudowany w mikrokontroler generator wzbudza drgania kryształu kwarcu, przez co układ może odmierzzać czas z wysoką dokładnością – co przecież jest jego głównym zadaniem. Kondensatory C7 i C8 ułatwiają wzbudzenie drgań Q1 oraz stabilizują jego pracę.

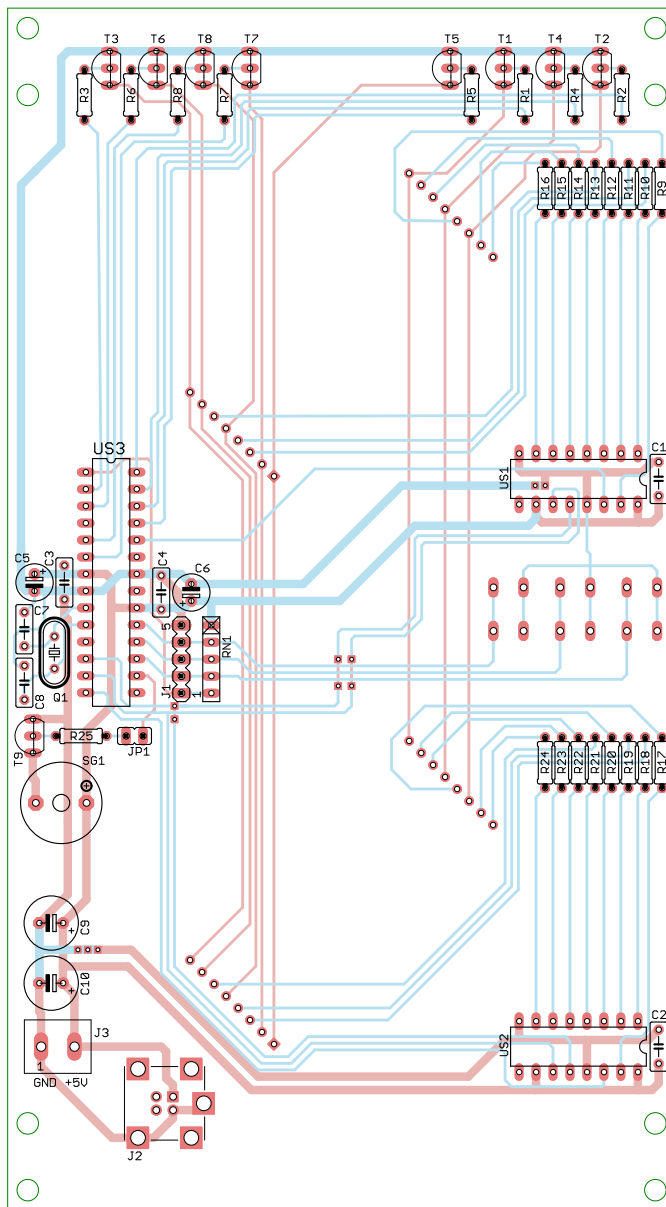
Użyte matryce LMD23088AUE-101A mają po 64 punkty świecące w kolorze czerwonym, ułożone w kwadrat 8×8. Anody diod są ułożone w kolumnach, natomiast katody w wierszach. Każdy wiersz jest załączany poprzez nasycenie jednego z tranzystorów NPN T1...T8 – w obu matrycach jednocześnie. Prąd zasilający wszystkie kolumny

wyływa z wyprowadzeń układów US1 oraz US2 i jest ograniczany przez rezystory o wartości 180 Ω. Zapewnia to dostatecznie wysoką jasność świecenia bez ryzyka uszkodzenia układów 74HC595. Wyświetlacze są odświeżane z częstotliwością 250 Hz; kolejne wiersze są załączane co 500 μs. Wiersze są wybierane międzyliniowo (1. wiersz – 3. wiersz – 5. wiersz – 7. wiersz – 2. wiersz itd.), a nie kolejnoliniowo, dla zmniejszenia efektu migotania wyświetlacza.

Użycie rejestrów przesuwanych było konieczne ze względu na niewystarczającą liczbę programowalnych



Rysunek 2. Schemat montażowy płytki



wyprowadzeń mikrokontrolera. Dane są wpisywane do obu rejestrów jednocześnie, ponieważ ich linie zegarowe są połączone, natomiast liniami danych mikrokontroler steruje oddzielnie. Dwukrotnie zmniejsza to czas potrzebny na aktualizację stanu świecenia punktów w wierszach.

Użytkownik ma do dyspozycji trzy przyciski monostabilne S1...S3, którymi steruje klepsydrą, oczym dalej. Rezystory podciągające, wbudowane w mikrokontroler, są wspomagane przez rezystory zawarte w drabince rezystorowej RN1, z których każdy ma rezystancję 10 kΩ. Zmniejsza to wrażliwość układu na zakłócenia elektromagnetyczne. Linia zerująca mikrokontrolera również została podciągnięta do dodatniego potencjału zasilającego w tym samym celu. Złącze J1 może służyć do zaprogramowania pamięci Flash mikrokontrolera US3 oraz do ustawienia jego bitów zabezpieczających.

Na płycie znalazło się gniazdo USB typu B, którym można zasiląć układ, jak również listwa

Wykaz elementów:

Rezystory: (THT o mocy 0,25 W)

R1...R8, R25: 1 kΩ
R9...R24: 180 Ω
RN1: 4×10 kΩ SIL9

Kondensatory:

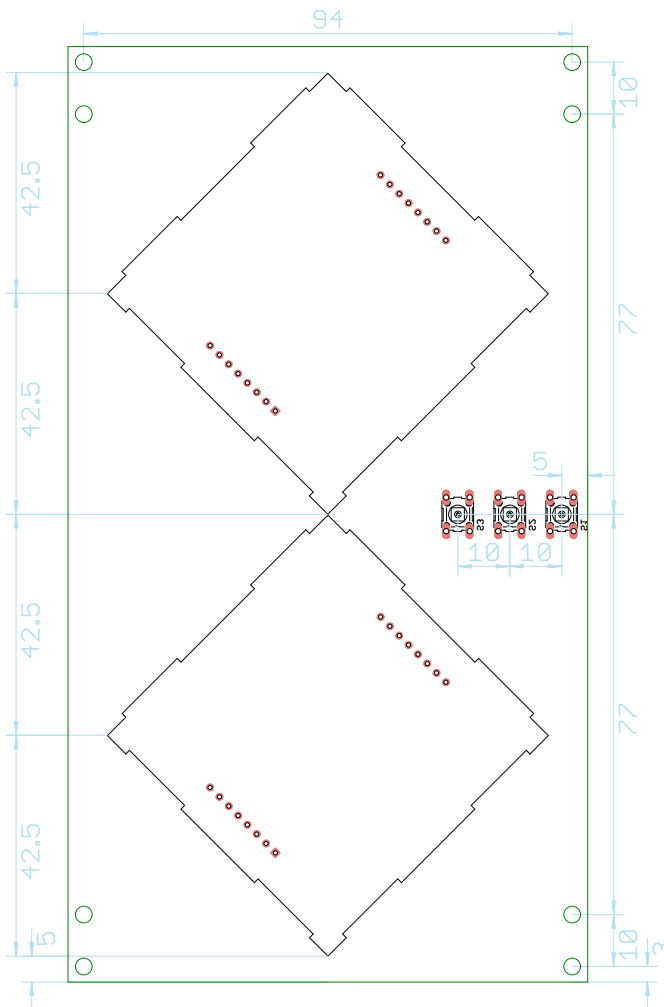
C1...C4: 100 nF raster 5 mm MKT
C5, C6: 22 μF/25 V raster 2,5 mm
C7, C8: 15 pF raster 5 mm monolityczne
C9, C10: 470 μF/16 V raster 3,5 mm

Półprzewodniki:

LED1, LED2: LMD23088AUE-101A
T1...T9: BC337 TO92
US1, US2: 74HC595 DIP16
US3: ATmega8A-PU DIP28

Pozostałe:

J1: goldpin 5 pin męski 2,54 mm THT
J2: UBBS-4R-D14-4D
J3: ARK2/500
JP1: goldpin 2 pin męski 2,54 mm THT + zworka
Q1: 16 MHz niski
S1...S3: microswitch 6×6 13,5 mm
SG1: PIEZO GEN 5 V
Jedna podstawka DIP28 wąska
Dwie podstawki DIP16



Rysunek 3. Rozmieszczenie przycisków, matryc LED oraz otworów montażowych na powierzchni płytki drukowanej

zaciśkowa. Pochodzące z zewnątrz napięcie stałe o wartości 5 V jest filtrowane przez łącznie osiem kondensatorów o zróżnicowanej pojemności, dla lepszego odsprężania zasilania dla mikrokontrolera oraz towarzyszących mu dwóch układów cyfrowych. Odświeżanie zawartości wyświetlaczy generuje znaczące tętnienia pobieranego prądu, co przekłada się na tętnienia zasilania, toteż konieczna była tak rozbudowana filtracja.

Montaż i uruchomienie

Układ został zmontowany na dwustronnej płytce drukowanej o wymiarach 180 mm × 100 mm. Jej wzór ścieżek oraz schemat montażowy przedstawia rysunek 2. Szczegóły na temat lokalizacji obiektów kluczowych dla wykonania obudowy (przycisków monostabilnych, matryc LED oraz ośmiu otworów montażowych) znajdują się na rysunku 3. Wszystkie otwory montażowe mają średnicę 3,2 mm.

Montaż proponuję rozpocząć od elementów o najmniejszej wysokości obudowy, czyli rezystorów. Pod układy scalone US1...US3 proponuję zastosować podstawki, aby ułatwić ich wymianę w razie uszkodzenia. W układzie prototypowym wyświetlacze LED1 i LED2 oraz przyciski S1...S3 znalazły się na wierzchniej stronie płytki (TOP), zaś cała reszta elementów na stronie spodniej (BOTTOM). Zmontowany układ można zobaczyć na **fotografii 1** oraz **fotografii 2**. Nic nie stoi na przeszkodzie, by przyciski wluutować od drugiej strony laminatu, podobnie sygnalizator dźwiękowy SG1.

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

Low Fuse = 0x3F

High Fuse = 0xD9

Szczegóły są widoczne na **rysunku 4**, który zawiera widok okna konfiguracji tychże bitów z programu BitBurner. W ten sposób zostanie uruchomiony wbudowany generator dla rezonatorów kwarcowych o wysokiej częstotliwości drgań oraz Brown-Out Detector, który wprowadzi mikrokontroler w stan zerowania, jeżeli jego napięcie zasilające spadnie poniżej 4 V. To znacznie zmniejsza ryzyko zawieszenia się mikrokontrolera podczas uruchamiania.

Poprawnie zaprogramowany układ jest gotowy do działania po podłączeniu zasilania do zacisków złącza J2 lub J3. Powinno to być napięcie stałe, dobrze filtrowane, najlepiej stabilizowane. Może pochodzić, na przykład, z ładowarki impulsowej ze złączem USB. Nominalna wartość tego napięcia powinna wynosić 5 V, lecz dopuszczalne są pewne odstępstwa. Maksymalna wartość to 5,5 V i wynika z ograniczeń nałożonych przez notę katalogową producenta mikrokontrolera, zaś dolną granicę można oszacować

REKLAMA

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

WWW.AKSOTRONIK.COM.PL

Aksotronik
ELEMENTY ELEKTRONICZNE

Magnesy neodymowe
oraz ferrytowe
Ceny od 0.10zł

Kostki elektryczne
zaciskowe
Ceny od 0.22zł

Szczotki węglowe
do elektronarzędzi
Ceny od 2.60zł/kpl

Pudełka/organizery
Ceny od 0.95zł

Przełączniki klawiszowe
wodoszczelne/pyłoszczelne
Ceny od 2.40zł

Druty oporowe
od 0.16 do 0.81mm
Ceny od 5.70zł

Przełączniki do elektronarzędzi
zwykłe i elektromagnetyczne
Ceny od 7.00zł

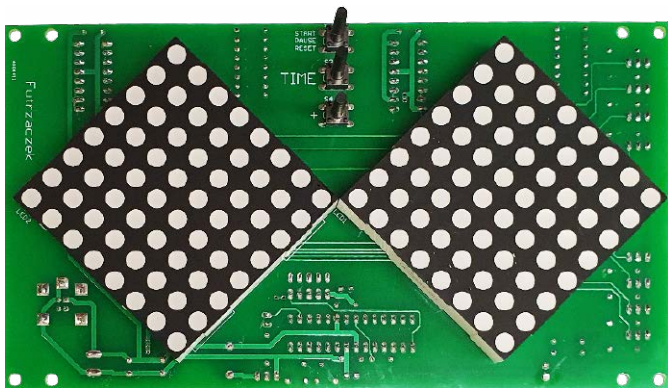
Zestawy śrubek M2, M3
z nakrętkami i podkładkami
Ceny od 2.50zł

Prowadniki
do przewodów
Ceny od 11.00zł

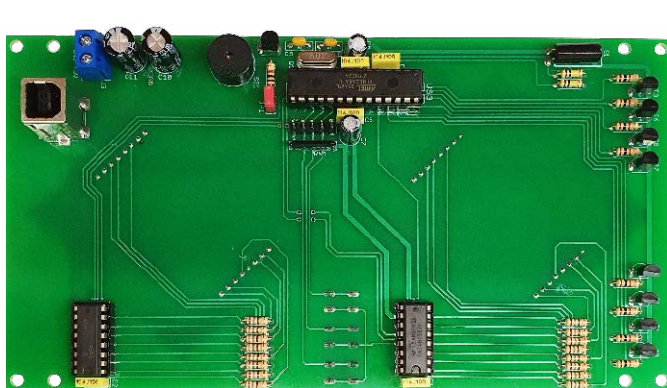
Złącza hermetyczne
Superseal
Ceny od 1.10zł/kpl

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



Fotografia 1. Widok zmontowanej płytki prototypowej – strona TOP



Fotografia 2. Widok zmontowanej płytki prototypowej – strona BOTTOM

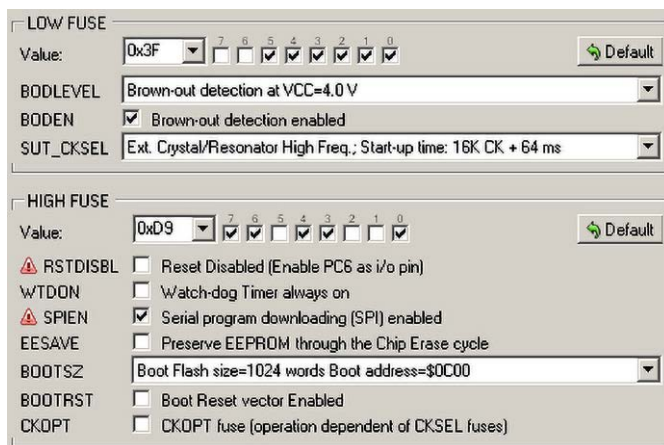
na 4,5 V – tak, by zabezpieczenie BOD było dalekie od zadziałania. Pobór prądu przez układ nie przekracza 120 mA przy zasilaniu go napięciem 5 V.

Rysunek 5 zawiera szczegółowe informacje dotyczące rozmieszczenia punktów świetlnych na powierzchni matryc LED. Ta informacja może być użyteczna do wykonania obudowy: jeżeli ktoś chciałby wykonać obudowę, w której każdy punkt byłby schowany za odrębnym otworem, te informacje wraz z rysunkiem 3 będą nieodzowne. Można jednak do sprawy podejść prościej i przesłonić obie matryce cienką, czerwoną płytą z przezrystego poliwęglanu, która spełni jednocześnie rolę filtra i zwiększy kontrast wyświetlaczy. Natomiast w wersji najprostszej proponuję przykręcenie dwóch małych kątowników do otworów montażowych na dole płytki, co umożliwi jej postawienie w pionie – po to są przewidziane po dwa otwory przy każdym narożniku, by ów kątownik stabilnie trzymał się laminatu. Oczywiście, można do tematu podejść również zupełnie inaczej, to są jedynie moje propozycje.

Eksplotacja

Po włączeniu zasilania, układ pokaże na matrycach ostatnio ustawiony czas z przedziału 1...99 minut, jak na fotografii 3. Górny wyświetlacz (LED1) wskazuje cyfrę dziesiątek, zaś dolny (LED2) cyfrę jedności. Jeżeli pamięć EEPROM mikrokontrolera była dotychczas pusta, domyślnym wskazaniem jest 1 minuta. W tym stanie układu można tę wartość zwiększać poprzez jednoczesne wciskanie przycisków S2 i S3 – po to, aby przypadkowe wciśnięcie jednego z nich nie przestawiło tej wartości. Ustawianie następuje w pętli: po doliczeniu do 99 minut, następnym krokiem jest 1 minuta. Dłuższe trzymanie tych przycisków powoduje przyspieszenie przewijania. Po ich zwolnieniu, zawartość nieulotnej pamięci EEPROM jest aktualizowana.

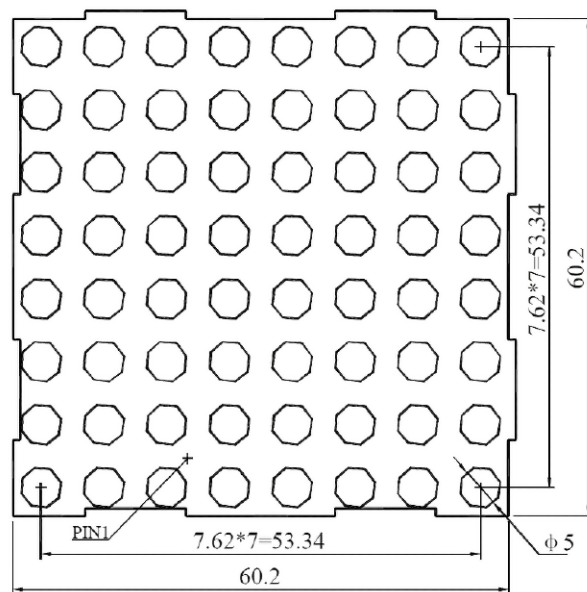
Z tego stanu można przejść do trybu odmierzania czasu, co czyni się poprzez krótkotrwałe wciśnięcie S1. Matryce będą wtedy wyglądały jak na fotografii 4 – górna połowka klepsydry będzie wypełniona punktami, zaś dolna zostanie



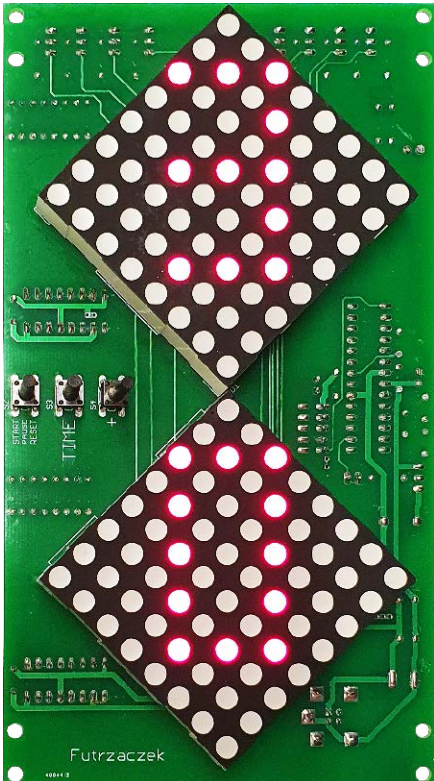
Rysunek 4. Szczegóły ustawienia bitów zabezpieczających

opróżniona. Wciskając S2, układ wróci do ustawiania czasu. Wciskając S1 jest raz, układ wyda z siebie pojedynczy, krótki pisk i rozpocznie odmierzanie ustawionego wcześniej czasu.

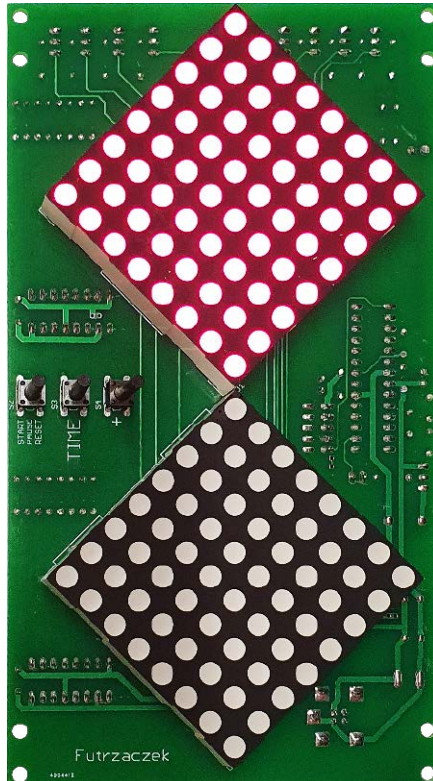
Kropki będą znikły z górnej matrycy i jednocześnie pojawiały się w dolnej, jak na fotografii 5. W tym



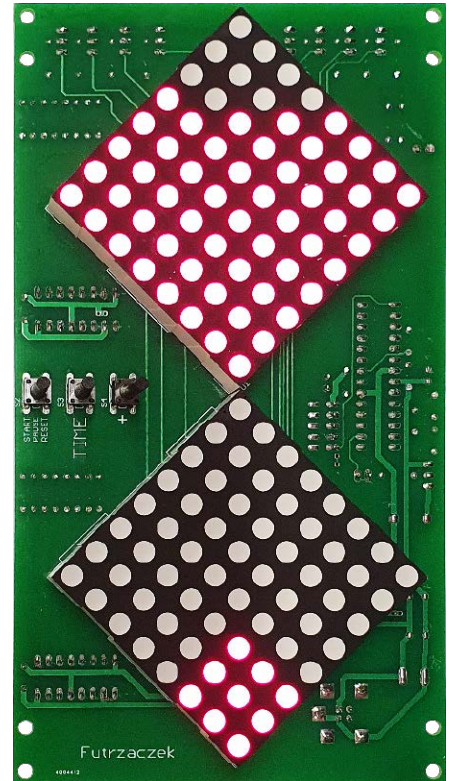
Rysunek 5. Szczegółowe wymiary wyświetlaczy matrycowych typu LMD23088AU-101A



Fotografia 3. Ustawianie czasu do odmierzenia



Fotografia 4. Stan początkowy trybu odmierzenia czasu



Fotografia 5. Wygląd matryc podczas odmierzenia czasu

stanie można wstrzymać odliczanie, wciskając na krótko S1, co zostanie oznajmione dwoma krótkimi piśnięciami. Obie matryce zaczną migać z częstotliwością 1 Hz, zaś stan obu z nich zostanie niezmieniony. Wciskając S1 na krótko, układ wyda z siebie jedno krótkie piśnięcie i będzie kontynuował odliczanie czasu. Z kolei przytrzymując S1 na jedną sekundę, układ wróci do stanu początkowego trybu odmierzenia czasu (jak na fotografii 4); można zacząć liczenie od nowa (S1) lub zmienić ustawienie (S2).

Tempo przeskakiwania kropek na dół zależy od ustawionego czasu, po rozpoczęciu odliczania układ dzieli zadany

czas na 64 jednakowe interwały. Kiedy już wszystkie kropki „przesypią się” na dół, odmierzenie czasu dobiegnie końca. Układ wyda z siebie trzy dłuższe piśnięcia i wróci samoczynnie do początkowego stanu odmierzenia czasu – jak na fotografii 4.

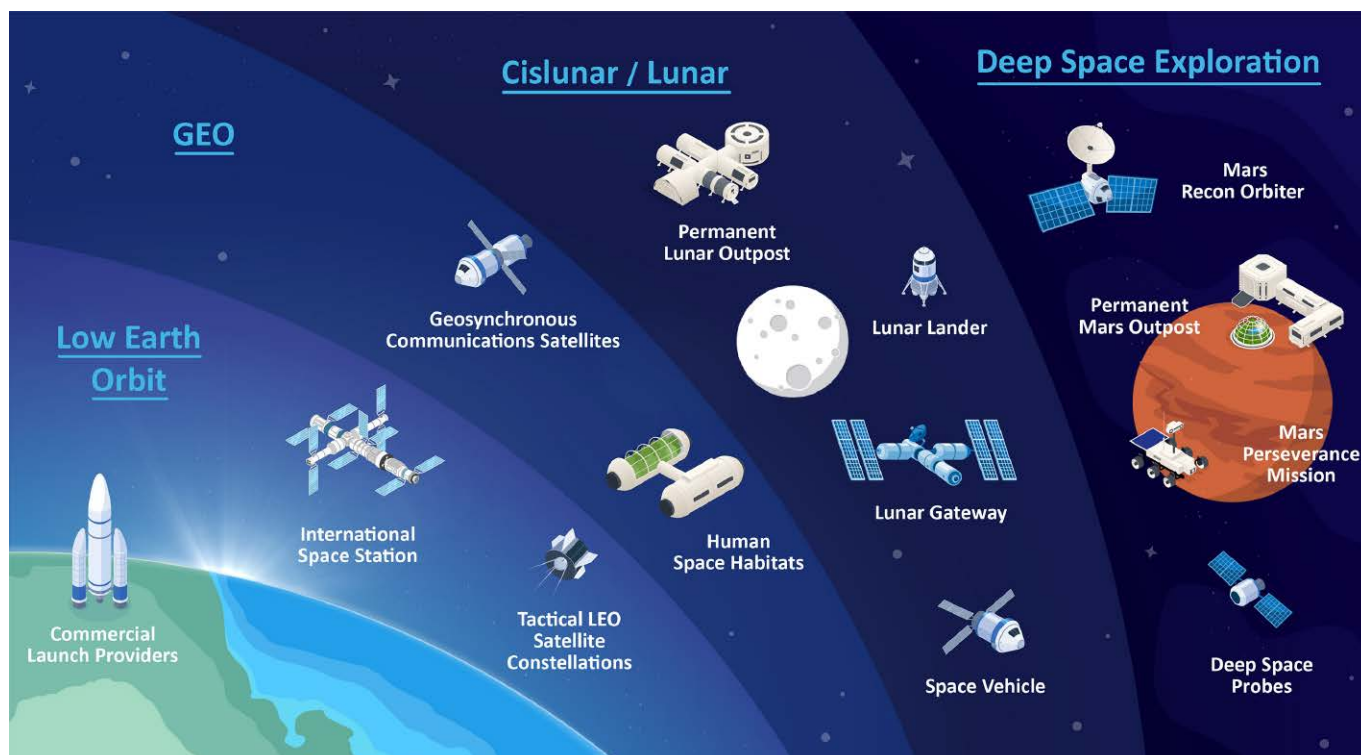
Sygnalizację dźwiękową układu można w bardzo łatwy sposób wyłączyć. Wystarczy zdjąć zworę JP1, aby tranzystor T9 nie mógł sterować buzzerem SG1. Wszystkie pozostałe zachowania układu pozostaną niezmienione.

Michał Kurzela, EP

m.technik
Ciekawi świata są zawsze młodzi

REKLAMA

przejrzyś i kupisz na stronie
www.ulubionykiosk.pl



Rysunek 1. Zastosowania półprzewodników w kolejnych obszarach przestrzeni kosmicznej – od niskiej orbity okołoziemskiej (LEO) i orbity geostacjonarnej (GEO), przez przestrzeń okotoksiężycową, po eksplorację dalekiego kosmosu

Jak rozwiązania półprzewodnikowe ewoluowały, by wspierać przemysł kosmiczny: od Apollo 11 do Starlink

Komponenty półprzewodnikowe to niedoceniani bohaterowie misji kosmicznych, zapewniający niezawodność i wydajność w ekstremalnym środowisku kosmosu. W ciągu ostatnich 60 lat firma Microchip odegrała kluczową rolę w ponad 100 misjach kosmicznych, przyczyniając się do sukcesu niektórych z najbardziej przełomowych kamieni milowych w eksploracji kosmosu. Od pierwszej udanej amerykańskiej misji kosmicznej w 1958 roku po trwające obecnie misje Artemis komponenty te niezmiennie potwierdzają swoją wartość.

Kluczowa rola komponentów półprzewodnikowych w misjach kosmicznych

Od czasu wyniesienia pierwszego amerykańskiego satelity, Explorera 1, na pokładzie rakiety Jupiter-C, podzespoły firmy Microchip wykazują kwalifikacje typu space heritage (dziedzictwa lotów kosmicznych), spełniając surowe normy odporności na promieniowanie i niezawodności podczas pracy w przestrzeni kosmicznej.

Rola komponentów półprzewodnikowych w misjach kosmicznych zaczęła się od układów kontroli częstotliwości, które miały kluczowe znaczenie we wczesnych

misjach. Układy te – takie jak oscylatory kwarcowe, oscylatory SAW sterowane napięciem (VCSO) czy zegary atomowe – są krytyczne w elektronice misji kosmicznych, ponieważ zapewniają dokładną transmisję i odbiór sygnałów, utrzymując stabilność łączności, integralność danych oraz synchronizację systemów. Komponenty te były niezbędne dla powodzenia pierwszej amerykańskiej misji kosmicznej w 1958 roku i położyły podwaliny pod dziedzictwo niezawodności w kosmosie. Lądowanie Apollo 11 na Księżycu w 1969 roku – jedno z największych osiągnięć ludzkości – również opierało się na tych technologiach. Microchip dostarczył krytyczne wsparcie łączności

w module księżycowym (LM) Apollo 11 na powierzchni Księżyca.

Oscylatory rubidowe, SAW i kwarcowe obsługują więcej zastosowań w komunikacji wojskowej, naziemnych stacjach satelitarnych oraz w aparaturze kontrolno-pomiarowej niż jakiegokolwiek inne precyzyjne źródła odniesienia częstotliwości na świecie.

Misja Voyager 1 – obecnie najdalej oddalony od Ziemi obiekt stworzony przez człowieka – dodatkowo pokazała niezrównaną wydajność półprzewodników w kosmosie. Elektronika Voyagera 1 łączyła układy scalone logiki TTL (transistor-transistor logic) oraz CMOS (complementary metal-oxide-semiconductor), komponenty analogowe, układy pamięci i dedykowane półprzewodniki zaprojektowane tak, by sprostać wyzwaniom dalekiego kosmosu. Główny komputer Voyagera 1 wykorzystywał dedykowany procesor (CPU) o nazwie SPS-8, zaprojektowany dla tej sondy przez NASA. Układy logiczne TTL były wówczas popularnym typem cyfrowych układów scalonych, podczas gdy obecnie większość półprzewodnikowych układów scalonych bazuje na technologii CMOS.

W ostatnich latach technologie półprzewodnikowe odgrywają centralną rolę w eksploracji Marsa. Łaziki Curiosity i Perseverance, które dostarczyły bezcennych informacji o Czerwonej Planecie, korzystają z tych komponentów, aby działać w trudnym środowisku Marsa. Łazik marsjański – w szczególności Perseverance – zawiera kilka komponentów firmy Microchip Technology. Należą do nich mikrokontrolery wykorzystywane w różnych systemach sterowania i zadaniach przetwarzania danych, a także układy zarządzania zasilaniem (power management IC), kluczowe dla efektywnego rozprowadzania energii do poszczególnych podzespołów łazika. Wszystkie te elementy są komponentami odpornymi na radiację (Rad-Hard), co gwarantuje, że elektronika wytrzyma trudne warunki kosmiczne. Komponenty te są niezbędne do działania łazika i pomagają mu wykonywać zadania naukowe na Marsie.

Jeśli chodzi o eksplorację Księżyca, misja Chandrayaan-3 – trzecia indyjska misja księżycowa – wykorzystała kilka komponentów półprzewodnikowych; podzespoły takie jak układy FPGA tolerujące promieniowanie (Radiation-Tolerant, RT) mają kluczowe znaczenie dla powodzenia misji, umożliwiając komunikację, nawigację i eksperymenty naukowe na powierzchni Księżyca.

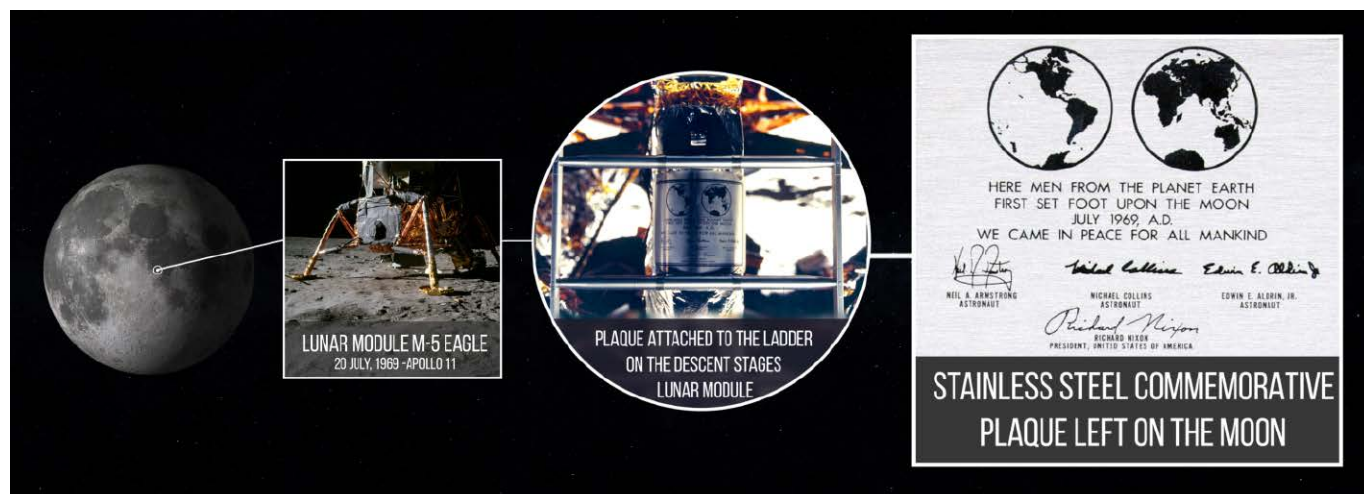
Trwające misje Artemis, których celem jest powrót człowieka na Księżyc, a docelowo wysłanie ludzi na Marsa, również opierają się na sprawdzonej wydajności i niezawodności technologii półprzewodnikowych.

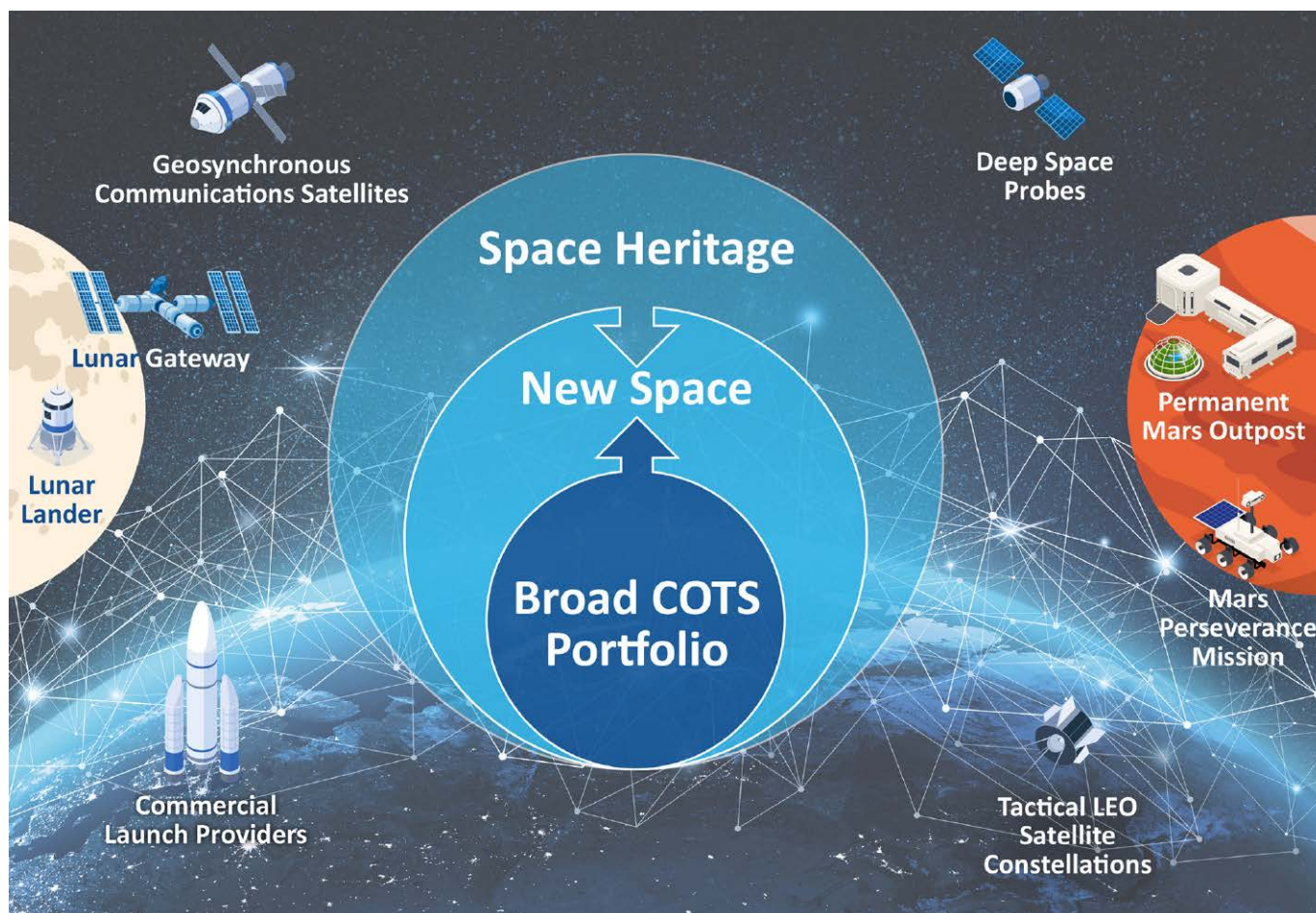
Znaczenie niezawodności i wydajności w programach kosmicznych

W trudnym środowisku kosmosu niezawodność i wydajność nie są jedynie ważne – są krytyczne. Komponenty półprzewodnikowe stanowią serce współczesnych misji kosmicznych, zasilając wszystko, od satelitów i łazików po systemy łączności i stacje kosmiczne. Z uwagi na ekstremalne warunki kosmosu – bezlitosne temperatury, intensywne promieniowanie i próżnię – komponenty muszą działać bezbłędnie przez długi czas. Nawet najmniejsza awaria półprzewodnika może doprowadzić do niepowodzenia misji, co podkreśla znaczenie doboru komponentów o najwyższej niezawodności.

Wyzwanie promieniowania w kosmosie

Kosmos jest wypełniony wysokim poziomem promieniowania, które może być niszczące dla komponentów elektronicznych. Promieniowanie to może degradować materiały, powodować awarie elektryczne i uszkadzać przesyłane dane. Na przykład środowisko promieniowania słonecznego poza ochronną atmosferą Ziemi może wystawiać komponenty na działanie cząstek o wysokiej energii, wywołujących zaburzenia od pojedynczych zdarzeń (single-event upsets, SEU) lub uszkodzenia od całkowitej dawki promieniowania.





Rysunek 2. Od dziedzictwa lotów kosmicznych (space heritage) i szerokiego portfolio COTS do segmentu new space – łączenie sprawdzonych technologii kosmicznych z komercyjnymi podzespołami

Aby sprostać tym wyzwaniom, zaawansowane techniki uodpornienia radiacyjnego obejmują stosowanie specjalistycznych materiałów, takich jak półprzewodniki odporne na promieniowanie, oraz modyfikacje konstrukcyjne ograniczające podatność na promieniowanie. Przykładowo mikroprocesory klasy kosmicznej stosowane w komputerach pokładowych są często uodporniane radiacyjnie już na poziomie projektu (radiation-hardened by design, RHBD), tak aby pojedyncza awaria nie wyłączyła całego systemu.

Wymagające testy symulujące warunki kosmiczne

Oprócz uodporniania radiacyjnego firmy z dziedzictwem w misjach kosmicznych opracowały rygorystyczne procesy testowania i kwalifikacji, które gwarantują niezawodność i wydajność ich komponentów. Testy te wykraczają daleko poza zwykłą kontrolę jakości produkcji. Półprzewodniki klasy kosmicznej przechodzą rozbudowane testy cyklicznego oddziaływania temperatury (thermal cycling), symulujące szerokie wahania temperatur w kosmosie – od intensywnego żaru Słońca po przejmujący chłód dalekiego kosmosu. Przykładem jest testowanie przez NASA komponentów łazika Mars Perseverance, które doświadczały wahań temperatury od -55°C do 125°C ,

co wymagało komponentów zdolnych wytrzymać takie ekstrema bez awarii.

Komponenty przechodzą również testy wibracyjne symulujące naprężenia i drgania występujące podczas startu. Ogromne siły generowane podczas startów rakiet są nieporównywalne z czymkolwiek, co spotykamy na Ziemi, dlatego komponenty półprzewodnikowe muszą wytrzymać te warunki bez utraty integralności. Na przykład podczas misji Apollo 11 krytyczna elektronika została poddana testom wibracyjnym, aby zapewnić jej przetrwanie potężnych sił startowych, co ostatecznie przyczyniło się do sukcesu lądowania na Księżycu.

Długoterminowa niezawodność: przykłady z historii kosmosu

Misje kosmiczne wymagają komponentów, które nie tylko działają w trakcie misji, ale również niezawodnie funkcjonują przez długi czas. Sonda Voyager 1, wystrzelona w 1977 roku, jest doskonałym przykładem tego, jak kluczowa jest niezawodność w przypadku misji długotrwałych. Przez ponad 40 lat w kosmosie sonda nadal komunikuje się z Ziemią – dzięki komponentom półprzewodnikowym, które uodporniono radiacyjnie i rygorystycznie przetestowano pod kątem odporności na ekstremalne warunki.

Innym przykładem jest Międzynarodowa Stacja Kosmiczna (ISS), która opiera się na licznych systemach półprzewodnikowych utrzymujących systemy podtrzymywania życia, prowadzących eksperymenty naukowe i zapewniających otwarte kanały łączności. ISS jest nieustannie narażona na trudne środowisko promieniowania kosmicznego, przy temperaturach sięgających od $+121^{\circ}\text{C}$ (po stronie nasłonecznionej) do -157°C (w cieniu) – a mimo to znajdujące się na pokładzie komponenty półprzewodnikowe muszą działać niezawodnie dzień po dniu.

Przyszłość komponentów półprzewodnikowych w kosmosie

W miarę jak przemysł kosmiczny wciąż się rozwija – wraz z rozkwitem konstelacji satelitów na niskiej orbicie okołoziemskiej (LEO) oraz postępującą komercjalizacją kosmosu – rośnie też zapotrzebowanie na niezawodne i wysokowydajne komponenty półprzewodnikowe. Nowe przedsięwzięcia kosmiczne wymagają komponentów zdolnych sprostać unikalnym wyzwaniom, łączącym wysoką niezawodność, innowacyjność i cele nastawione na zysk.

Poszerzanie granic: ewolucja rynku kosmicznego

Rosnącym trendem w przemyśle kosmicznym jest coraz częstsze wykorzystanie podzespołów komercyjnych z półki (Commercial Off-The-Shelf, COTS), które dzięki natychmiastowej dostępności stanowią opłacalne rozwiązanie dla misji kosmicznych. Sieć satelitarna Starlink wykorzystuje komponenty COTS, aby obniżyć koszty i przyspieszyć produkcję. Wiele układów elektroniki pokładowej – zwłaszcza w systemach niekrytycznych – opiera się na tych komponentach, zarówno przystępnych cenowo, jak i starannie dobranych pod kątem niezawodności w środowisku kosmicznym.

Dobrym przykładem jest europejska rakieta nośna Ariane, opracowana wspólnie z ESA (Europejską Agencją Kosmiczną) i CNES (francuską agencją kosmiczną). Ariane 5 (1985) była wyposażona w uodporniony radiacyjnie procesor centralny Sparc w klasie QML, w hermetycznej obudowie, oraz w magistralę 1553 do komunikacji między wszystkimi systemami rakiety. Najnowsza wersja, Ariane 6 (2024), wykorzystuje obecnie procesor COTS oparty na architekturze Arm®, z montażem w obudowach z tworzywa sztucznego oraz magistralą Ethernet do komunikacji – będącą również szeroko przyjętym standardem przemysłowym, w przeciwieństwie do technologii typowo kosmiczno-wojskowej stosowanej w Ariane 5.

W przypadku systemów o wysokiej niezawodności, które muszą utrzymać wysoką wydajność w łączności kosmicznej, urządzenia COTS wymagają jednak adaptacji i kwalifikacji, co wiąże się z koniecznością posiadania specjalistycznej wiedzy. Firmy chcące wejść na rynek nowej

kosmonautyki (new space) lub przejść z segmentu new space do dalekiego kosmosu (deep space) muszą współpracować z producentami półprzewodników o udokumentowanym dziedzictwie lotów, zdolnymi do podniesienia urządzeń COTS do poziomu spełniającego surowe wymagania misji kosmicznych.

Jak dziedzictwo lotów kosmicznych otwiera drogę do nowej ery dostępności i współpracy w zastosowaniach kosmicznych

Microchip oferuje klientom prostą ścieżkę przejścia od urządzeń COTS do układów kwalifikowanych do zastosowań kosmicznych. Podejście to zapewnia skalowalne i konfigurowalne rozwiązania dopasowane do unikalnych wymagań każdej misji. Taka elastyczność jest niezbędna w szybko zmieniającym się i nieustannie ewoluującym przemyśle kosmicznym, w którym wciąż pojawiają się nowe technologie i profile misji. Obniżając bariery dla komercjalizacji i eksploracji kosmosu, Microchip umożliwia szerszy dostęp do technologii i innowacji kosmicznych.

W krótkiej perspektywie przyszłość półprzewodników w kosmosie można postrzegać jako połączenie różnorodnych strategii: podnoszenia klasy urządzeń COTS, wykorzystania doświadczenia w lotach kosmicznych poprzez wersje produktów w klasie sub-QML w celu ograniczenia wymagań dotyczących selekcji (screening), obniżenia kosztów i skrócenia czasu realizacji dostaw, a także dostosowywania procesów produkcyjnych do unikalnych wymagań konkretnych profili misji.

Łącząc te podejścia, branża będzie nadal zapewniać solidną i elastyczną przyszłość półprzewodników w kosmosie.

Podsumowanie

Komponenty półprzewodnikowe mają kluczowe znaczenie dla powodzenia misji kosmicznych, zapewniając niezbędną niezawodność i wydajność w trudnym środowisku kosmosu.

W miarę dalszego rozwoju przemysłu kosmicznego zapotrzebowanie na niezawodne i wysokowydajne komponenty półprzewodnikowe będzie tylko rosnąć.

Nicolas Ganry

starszy menedżer ds. marketingu w jednostce biznesowej lotnictwa, kosmonautyki i obronności (aerospace & defense) firmy Microchip Technology





M5Stack Cardputer ADV z modułem LoRa-1262

Kieszonkowy komputer (nie tylko) edukacyjny na ESP32-S3

Zestawów rozwojowych i edukacyjnych na rynku nie brakuje. Większość przyjmuje formę prostej płytki drukowanej z kilkoma złączami, paroma przyciskami i diodami LED, czasem z dodatkowymi sensorami lub/i wyświetlaczem. Zestawy takie często są bardzo tanie, ale przez to działają tylko w pracowni elektronika. Firma M5Stack przygotowała zestaw, który pracownię może opuścić, a przez to staje się praktycznym urządzeniem codziennego użytku.

Cardputer ADV, bo o tym zestawie mowa, to drugi kieszonkowy komputer w ofercie M5Stack. Zestaw ten oparty jest o moduł, również produkowany przez M5Stack, o nazwie Stamp S3A, którego sercem jest układ ESP32-S3FN8 firmy Espressif. Firma M5Stack oferuje szereg innych zestawów, skierowanych głównie do użytkowników sieci Meshtastic, z tego też powodu można zakupić Cardputer z dodatkowym modułem (Cap, czyli „czapka”)

LoRa 1262. Taki zestaw, z aktywną aplikacją Meshtastic przedstawia fotografia do tego artykułu.

Zakup, opis i specyfikacja

Jak wspomniano wyżej, Cardputer ADV jest napędzany układem ESP32-S3FN8. Jest to układ typu SoC ze zintegrowanym radiem 2,4 GHz dla łączności Wi-Fi i Bluetooth Low Energy. Sercem układu jest dwurdzeniowy mikroprocesor

oparty o architekturę Xtensa® 32-bit LX7. Częstotliwość taktowania wynosi maksymalnie 240 MHz. Ten wariant układu ESP32-S3 ma też zintegrowaną pamięć flash o pojemności 8 MB. Teoretycznie układ ten wspiera też zewnętrzną pamięć flash oraz RAM do 1 GB, choć sam posiada tylko 512 kB pamięci SRAM. Moduł Stamp S3A nie zawiera jednak żadnej dodatkowej pamięci zewnętrznej. Sam Cardputer ADV dodaje przede wszystkim ekran IPS o przekątnej 1,18 cala, 65-przyciskową klawiaturę, dwa gniazda rozszerzeń, gniazdo kart microSD, oraz akumulator. Wewnątrz kryją się też układ IMU, mikrofon oraz wzmacniacz z głośnikiem oraz dioda podczerwieni. Producent nie pokusił się o instalację kompletnego modułu do zapomnianej już obecnie łączności IrDA, co mogłoby poszerzyć zastosowanie tego małego, kieszonkowego komputera. Tak, Cardputer ADV można nazwać pełnoprawnym komputerem, gdyż pozwala na łatwe tworzenie i uruchamianie aplikacji, a nawet posiada kilka gier. Komputer w zestawie z modułem LoRa 1262 można zakupić na stronie AliExpress, i kosztuje w chwili pisania artykułu około 263 zł, z możliwością uzyskania zniżki 26 zł. Jest to uczciwa cena w zamian za możliwości i generalną użyteczność tego zestawu. Cardputer ADV bez modułu LoRa 1262 kosztuje około 175 zł.

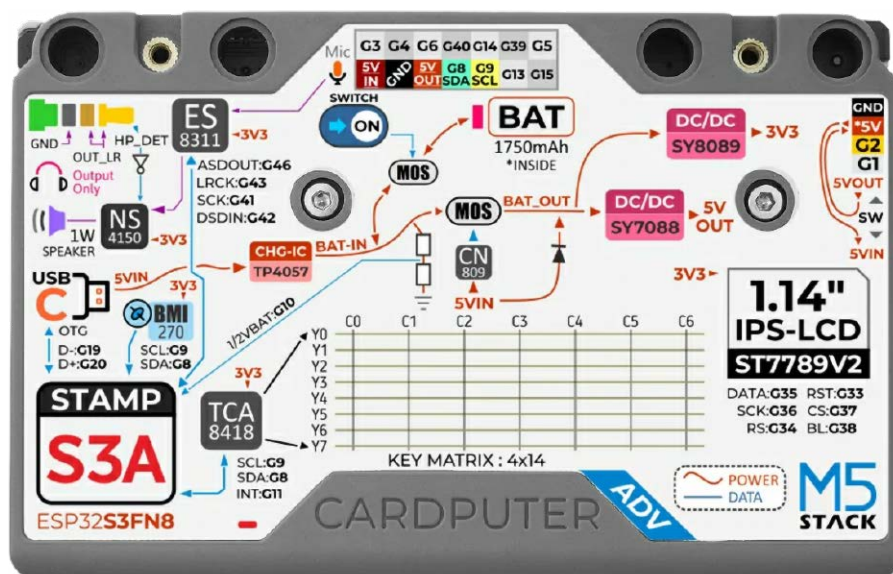
Urządzenie jest relatywnie małe i faktycznie ma rozmiar karty kredytowej, choć nie jej grubość. Moduł LoRa nieco zwiększa jego wymiary, a sposób montażu z jednej strony jest mało wygodny przy trzymaniu urządzenia w kieszeni, z drugiej sprawia, iż wygodniej trzyma się je w ręku. Moduł LoRa 1262 można przykręcić do Cardputera ADV za pomocą dwóch wkrętów walcowych M2×6 mm. Wkrętów brak w zestawie. W zestawie brak też karty microSD, warto jednak się w nią zaopatrzyć. Klawiatura jest doprawdy mała i dość niewygodna – producent zdecydował się użyć mikroprzełączników tact i obciął dodatkowo koszty nie dodając żadnych gumek. Autor używał podobnej wielkości klawiatur, na przykład w smartfonach HTC (na przykład model S740), które to wykonane były z gumowej membrany przycisków i metalowych kopulek, czyniąc je znacznie wygodniejszymi bez

znaczącego wzrostu kosztów produkcji. Klawiatura jest najłabszą stroną tego komputera, szczególnie gdy trzeba korzystać z kombinacji klawiszy – jest to zwyczajnie niewygodne. Klik jest twardy, a klawisze nadzwyczaj małe i ciasno upakowane.

Domyślnie komputer przychodzi z zainstalowaną aplikacją demonstracyjną (**fotografia 1**), która pozwala przetestować różne funkcje urządzenia. Domyślnie wybrane jest demo „Recorder”, które pozwala sprawdzić działanie mikrofonu, a także głośnika przez odtworzenie ostatnich dwóch sekund zarejestrowanego dźwięku. Innym, ciekawym demem jest Chat LoRa, który do przetestowania wymaga jednak dwóch zestawów Cardputer ADV z modułami LoRa 1262. Aplikacja demonstracyjna ma też na celu pokazanie możliwości SDK i biblioteki graficznej M5Stack UIFlow 2.0. Środowisko to dostępne jest w dwóch formach: aplikacji dla komputerów PC/Mac oraz w formie aplikacji



Fotografia 1. M5Stack Cardputer ADV z domyślną aplikacją demonstracyjną



Fotografia 2. Na spodniej stronie Cardputer ADV ma uproszczony opis wewnętrznego układu połączeń oraz opis pinów złączy

na Cardputer ADV. Dostępna też jest wersja online do użycia w przeglądarce. UIFlow 2.0 pozwala tworzyć aplikacje w sposób wizualny, przez składanie ich z „klocków”, ale udostępnia też język MicroPython. Do instalacji aplikacji służy narzędzie M5Burner, które zawiera kolekcję oficjalnych i nieoficjalnych aplikacji dla różnych zestawów M5Stack. Najlepszą metodą jednak jest zainstalowanie aplikacji M5Launcher, która pozwala instalować aplikacje z karty microSD. Inne aplikacje wtedy pobiera się narzędziem M5Burner, znajdując się w katalogu /packages/firmware. Zaleca się pobierać jedną aplikację naraz, a następnie stworzenie jej kopii i zmianę nazwy kopii na nazwę aplikacji. Pobrane aplikacje umieszcza się na karcie microSD w głównym katalogu.

Pewną ciekawostką jest umieszczenie na spodzie urządzenia nalepki z opisem wewnętrznej budowy i struktury urządzenia (fotografia 2), podobna „mapka” jest też na module LoRa 1262 (fotografia 3). W praktyce najistotniejszą informacją jest „pinologia” gniazd rozszerzeń. Warto dodać, iż M5Stack (i nie tylko) oferuje kompatybilne z nimi moduły czujników i innych rozszerzeń. Wbudowany akumulator ma pojemność 1750 mAh, a jego ładowanie odbywa się poprzez włączenie urządzenia i podłączenie kabla USB-C do gniazda w module Stamp S3A. Urządzenie można też zasilac przez czteropinowe złącze rozszerzeń z boku, po wcześniejszym przełączeniu kierunku pracy pinu +5 V tego złącza.

Praca z Cardputer ADV oraz aplikacje

Wspomniano już o kilku aplikacjach dla Cardputer ADV: Meshtastic, Demo, czy UIFlow 2.0. Ciekawą aplikacją jest WebRadio by WuSiU – proste radio internetowe Wi-Fi, dla którego listę stacji można zapisać na karcie microSD w prostym pliku tekstowym. Po skonfigurowaniu połączenia Wi-Fi aplikacja startuje od razu i pozwala wybrać stację z listy zapisanych adresów. Jakość dźwięku z wbudowanego głośnika jest adekwatna. Niestety, sygnał wyjściowy, również na wyjściu słuchawkowym, jest monofoniczny, i wynika to z konstrukcji urządzenia. Inny projekt radia internetowego opartego o moduł z ESP32, YoRadio, uzyskuje znacznie lepszą jakość dźwięku w stereo dzięki użyciu lepszych układów DAC do zastosowań audio. Cardputer ADV ma zamiast tego układ ES8311, budżetowy układ o przyzwoitych parametrach przeznaczony raczej do urządzeń IoT i zabawek, a nie dla sprzętu Hi-Fi. Wśród aplikacji dostępnych jest też kilka odtwarzaczy MP3/FLAC,



Fotografia 3. Moduł LoRa-1262 ma również ilustrację opisującą jego budowę

ale wygląda na to, iż w chwili pisania artykułu nie ma ani jednej aplikacji audio wspierającej użycie zewnętrznych słuchawek Bluetooth, mimo iż układ ESP32-S3FN8 wspiera Bluetooth Low Energy w wersji 5.0. Wersja ta wspiera przesyłanie strumieni audio, ale dopiero wersja 5.2 dodaje pełne, standardowe wsparcie z użyciem kodeka LC3.

Największą grupę aplikacji stanowią narzędzia do testowania i łamania zabezpieczeń sieci Wi-Fi czy urządzeń Bluetooth. Niektóre wspierają dodatkowo moduły radiowe na pasma sub-GHz oparte o układ CC1101. Widać tu wyraźną inspirację narzędziem Flipper Zero, które potrafi zdziałać więcej dzięki lepszemu zestawowi peryferiów. Flipper Zero stał się popularny dzięki serii wiralowych filmów na platformach społecznościowych, które sugerowały (błędnie), iż można nim „hakować” wszystko bez żadnej wiedzy. W rzeczywistości jest to użyteczne narzędzie dla ludzi posiadających zaawansowaną wiedzę na temat zabezpieczeń. Na tym tle aplikacje „hakerskie” dla Cardputer ADV należy traktować raczej w kategorii ciekawostki, zresztą połowa z nich to skanery Wi-Fi, czyli narzędzia, które zastąpić może prosta aplikacja na smartfona.

Na platformę dostępnych jest też kilkanaście gier i emulatorów, choć większość z nich wymaga zewnętrznego ekranu, gdyż dostępny w urządzeniu jest zwyczajnie za mały i ma złe proporcje dla emulacji takich systemów, jak na przykład GameBoy. Dodatkowo na rynku nie brakuje kieszonkowych konsolek emulujących systemy od automatów arcade i konsol z ostatnich trzech dekad XX wieku po systemy domowe i przenośne z początku XXI wieku. Jedyną kategorią, w której Cardputer ADV byłby dobrym wyborem są tekstowe gry RPG/eksploracyjne, jak na przykład Colossal Cave Adventure czy gry paragrafowe (Interactive Fiction – IF). Niestety, pod tym względem platforma nie ma zbyt wiele do zaoferowania, a szkoda. Nie

ma ani jednego projektu stworzenia portu maszyny wirtualnej „Z-Machine” czy interpretera języka Inform. Jedyną metodą by zagrać na przykład w grę Zork na Cardputer ADV jest uruchomienie emulatora procesora Zilog Z80 z systemem CP/M i uruchomienie na nim oryginalnej wersji gry Zork od firmy Infocom, która to stworzyła Z-Machine w 1979 roku by móc portować swoje gry na różne platformy.

Wśród aplikacji znaleźć można też kilka automatów perkusyjnych, syntezyatorów i przynajmniej jedno wirtualne pianino. Nieźle, biorąc pod uwagę ograniczenia sprzętowe platformy, szczególnie klawiaturę. Jest też kilka aplikacji dla modułu GPS, lecz Autor miał problemy z uzyskaniem danych o swojej pozycji. Aplikacja do obsługi Meshtastic korzysta z GPS by pozyskać aktualny czas, a układ ESP32-S3FN8 ma wbudowany zegar czasu rzeczywistego. Niestety, bez dedykowanego pinu VBATT nie ma możliwości dodania niezależnego podtrzymania tego peryferium, a M5Stack nie dla oszczędności nie dodało niezależnego układu RTCC z własnym podtrzymaniem. Wśród aplikacji znajdują się przynajmniej dwie „klawiatury” Bluetooth i jedna „klawiatura/myszka” Bluetooth/USB. Klient SSH też jest dostępny, choć wizja używania go na tak małym ekranie jest mało atrakcyjna, biorąc pod uwagę, iż nie da się na nim zmieścić nawet czterdziestu kolumn, nie wspominając o osiemdziesięciu. Nie brakuje też pilotów zdalnego sterowania opartych o diodę IR wbudowaną w Cardputer ADV. Ostatnią grupą aplikacji są asystenci/chatboty AI, które robią to samo, co aplikacje dostępne na smartfony, tylko gorzej, bo Cardputer ADV nie ma kamery. Nawet nie warto sobie nimi zaśmiecać karty pamięci.

Po przejrzeniu dość długiej listy aplikacji, na której wiele pozycji się powtarzało należy dojść do wniosku, iż nie ma zbyt wielu naprawdę użytecznych rzeczy dla tej platformy. Aplikacje emulujące pilot zdalnego sterowania, WebRadio by WuSiU oraz aplikacje do sieci Meshtastic są najbardziej użyteczne, obok M5Launcher do ładowania ich z karty pamięci. Wszystkie aplikacje „hakerskie” to zasadniczo ograniczone możliwościami sprzętu ciekawostki, często będące kopiami, portami lub klonami innych aplikacji. Wbrew pozorom jednak obecność dziesiątek mało użytecznych aplikacji jest dobrym znakiem, gdyż pokazuje, jak łatwo tworzy się i portuje oprogramowanie na tę platformę. Warto pamiętać o dużo większym zbiorze gotowych aplikacji i przykładów dla całej platformy ESP32, nie tylko na tym, co oferuje firma M5Stack i użytkownicy jej platformy. Stopień integracji różnych elementów w Cardputer ADV czyni go też dobrym punktem wyjścia do samodzielnej nauki tworzenia aplikacji, zwłaszcza iż M5Stack stara się to zadanie ułatwić.

Od strony ergonomii Cardputer ADV sprawdza się wcale nieźle. Wrażenie psuje klawiatura, a i ekran może być dla

niektórych zbyt mały. Moduł LoRa-1262 nie przeszkadza, a wręcz ułatwia trzymanie urządzenia. By wgrać nowe oprogramowanie można skorzystać, jak wspomniano wcześniej, z aplikacji M5Burner – w tym celu należy przytrzymać przycisk Go przed podłączeniem urządzenia do komputera by włączyć tryb bootloadera. Przycisk ten jest też zdublowany na module Stamp S3A. Preferowaną metodą jednak jest użycie karty pamięci na firmware i wgranie aplikacji M5Launcher.

Programowanie Cardputer ADV może odbywać się na dwa sposoby. Pierwszym jest UIFlow 2.0 Web IDE – rozwiązanie w sam raz dla celów edukacyjnych. Oferuje dwie metody programowania: za pomocą bloków, z których składa się program, lub w języku microPython. Wadą jest konieczność korzystania ze środowiska w chmurze, w dodatku z obowiązkiem rejestrowania się na platformie M5Stack. Dostępne jest środowisko UIFlow Desktop IDE, ale w wersji 1.0, niekompatybilnej z Cardputerem, bawet w starszej, bazowej wersji urządzenia. Dla pracy lokalnej do wyboru jest Arduino IDE albo VS Code z rozszerzeniem PlatformIO. W przypadku Arduino IDE należy ręcznie zainstalować pakiet wsparcia dla ESP32, a następnie biblioteki dla płytek ESP32 i oddzielnie biblioteki dla płytek M5Stack. Na liście będzie dostępna płytka Cardputer, ale jest ona kompatybilna z Cardputer ADV. W przypadku PlatformIO wystarczy wybrać ESP32-S3, a następnie trzeba dodać biblioteki M5Unified lub M5Cardputer. Zaawansowany użytkownik może skorzystać z ESP IDE od firmy Espressif, zyskując pełnię kontroli nad sprzętem, lecz w zamian musi przygotować własne sterowniki i skonfigurować piny by uzyskać dostęp do możliwości Cardputera.

Podsumowanie

M5Stack tworząc Cardputer ADV przygotował świetny zestaw edukacyjny, który nie wymaga niczego więcej do bycia użytecznym urządzeniem. Pozwala on na realizację projektów nie wymagających plątania przewodów i modułów czy przywiązania do zasilacza. Klawiatura może i jest niewygodna, ale jakość wykonania reszty oraz integracja wszystkiego, co potrzebne aż nadto to rekompensują. Moduł LoRa-1262 rozszerza funkcjonalność całości, a sposób jego montażu poprawia ergonomię. Taki zestaw może zainteresować elektroników chcących zapoznać się z siecią Meshtastic. Niestety, sieć ta występuje jedynie w dużych miastach, ale to natura infrastruktury tworzonej oddolnie przez hobbyistów. Cardputer ADV z modulem LoRa-1262 może stanowić wzór dobrej integracji komponentów dla innych producentów.

Paweł Kowalczyk

AI wbudowane we współczesne urządzenia elektroniczne

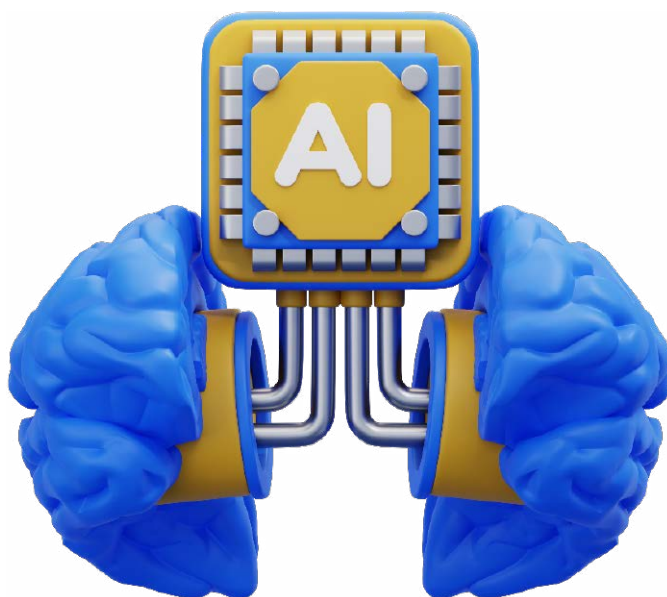
Sprzęt, który musiał dorosnąć

Jeszcze dziesięć lat temu uruchomienie modelu sieci neuronowej na mikrokontrolerze brzmiało jak żart. Dziś jest codzienną praktyką. Nie dlatego, że sieci neuronowe stały się prostsze – dlatego, że sprzęt przeszedł głęboką metamorfozę, której tempo zaskoczyło nawet samych projektantów układów scalonych.

Trzy ery osadzania inteligencji w sprzęcie

Huang i współpracownicy w swoim przeglądzie Embedded Artificial Intelligence opublikowanym w 2025 roku w czasopiśmie Electronics wyróżniają trzy kolejne ery, które ukształtowały obecny krajobraz EAI. Ta periodyzacja jest przydatna nie jako akademicka ciekawostka, lecz dlatego, że wyjaśnia, dlaczego konkretne narzędzia i podejścia, z którymi inżynier spotyka się dziś, wyglądają tak, a nie inaczej.

Era embrionalna, trwająca do 2016 roku, to czas pierwszych prób miniaturyzacji modeli chmurowych. Dominowało pytanie: jak sprawić, żeby sieć neuronowa była mniejsza? Odpowiedzią były wczesne algorytmy kompresji, takie jak SqueezeNet, który zastępował typowe filtry konwolucyjne strukturami znacznie bardziej zwartymi. Sprzęt nie nadążał – mikrokontrolery z tamtego okresu po prostu nie miały zasobów, żeby uruchomić nawet skompresowany model w czasie rzeczywistym.



Eksploracja mobile AI w latach 2017-2019 przyniosła przełom na obu frontach jednocześnie: architekturę MobileNet – sieć zaprojektowaną od zera pod kątem urządzeń mobilnych – oraz TensorFlow Lite, środowisko uruchomieniowe umożliwiające inferencję na smartfonach i tabletach. Urządzenia mobilne stały się platformą testową dla idei, które wkrótce miały zejść jeszcze niżej – na mikrokontrolery.

Era TinyML, trwająca od 2020 roku do dziś, to radykalne zejście po drabinie zasobów: z procesorów aplikacyjnych smartfonów na mikrokontrolery z pamięcią liczoną w kilobajtach. Kluczem były dwie równoległe zmiany: pojawienie



Rysunek 1. Trzy ery Embedded AI: od wczesnych algorytmów kompresji (SqueezeNet) przez MobileNet i TFLite do TinyML na mikrokontrolerach z akceleratorami AI i inferencją na poziomie mikrowatów

się mikrokontrolerów z wbudowanymi akceleratorami AI (Neural Processing Units – NPU, lub instrukcjami wektorowymi SIMD) oraz opracowanie modeli zdolnych do pracy przy poborze mocy na poziomie mikrowatów. Wynikiem jest kategoria urządzeń określana jako always-on: systemy, które nieprzerwanie monitorują otoczenie, reagują na zdarzenia i podejmują decyzje – a ich bateria wystarcza na miesiące lub lata.

Cztery klasy sprzętu – kiedy sięgać po którą?

Nie istnieje jeden „procesor do Edge AI”. Projektant ma do wyboru cztery zasadnicze klasy sprzętu, każda z innym profilem możliwości, ograniczeń i kosztów. Wybór między nimi nie jest decyzją techniczną – jest decyzją systemową, podejmowaną równoległe z projektowaniem architektury modelu ML.

Mikrokontroler pod TinyML: twardy świat liczb

Inżynier, który po raz pierwszy słyszy: uruchom model ML na mikrokontrolerze, musi zmierzyć się z konkretną arytmetyką. Heydari i Mahmoud, analizując w 2025 roku

Przyszłość EAI zmierza ku kolejnemu przełomowi paradygmatu: od statycznej inferencji do dynamicznego, adaptacyjnego uczenia.

– Huang X. et al.
Electronics 2025, 14(17), 3468

szereg eksperymentów TinyML z literatury naukowej, opisują typowy profil sprzętowy platform używanych w produkcyjnych wdrożeniach. Obraz jest jednoznaczny: działamy w środowisku ekstremalnie ograniczonym.

ARM Cortex-M4 jest dziś najczęściej spotykaną platformą referencyjną w badaniach TinyML. Jego popularność wynika z dostępności sprzętowej jednostki zmiennoprzecinkowej (FPU) i biblioteki CMSIS-DSP, zoptymalizowanej pod kątem operacji sygnałowych i macierzowych. Cortex-M7 oferuje dwukrotnie wyższą szybkość przetwarzania dzięki 6-etapowemu potokowaniu, ale przy wyższym poborze mocy – wybór między nimi to typowy kompromis dokładność/energia.

Nową klasę otwiera architektura Armv8.1-M z procesorem wektorowym Helium (MVE). Instrukcje SIMD Helium pozwalają na równoległe przetwarzanie do 16 operacji zmiennoprzecinkowych w jednym cyklu,

Tabela 3. Cztery klasy sprzętu Edge AI – profile możliwości i zastosowań. Opracowanie własne na podst.: Heydari i Mahmoud 2025; Huang et al. 2025; Abou Ali i Dornaika 2025

MCU (Mikrokontroler)	FPGA (Programowalna matryca)
Taktowanie: 100...480 MHz Flash: poniżej 1 MB SRAM: 128...512 kB Pobór mocy: 25...300 mW Koszt: kilka-kilkanaście USD Kiedy używać: → always-on sensing → wearables i IoT nodes → predictive maintenance → klasyfikacja audio/gestów	Taktowanie: 100...700 MHz Pamięć: BRAM wbudowany Moc: 1...25 W (zależnie od rozmiaru) Koszt: dziesiątki-setki USD Kiedy używać: → niskie serie produkcyjne → potrzeba rekonfigurowalności → prototypowanie ASIC → przetwarzanie w dziedzinie czasu
ASIC/NPU (Dedykowany układ)	SoC (System on Chip)
Wydajność: do kilkudziesięciu TOPS Moc: 0,3...15 W Koszt: kilka USD (masowa produkcja) Czas wdrożenia: długi (NRE) Kiedy używać: → masowa produkcja (miliony szt.) → znany, stały model ML → max wydajność na wat → wymagania certyfikacyjne	CPU: 1...8 rdzeni (Cortex-A/RISC-V) Akcelerator AI: NPU/DSP/GPU Pamięć: GB DRAM Moc: 2...30 W Kiedy używać: → złożona wizja komputerowa → wielomodalne AI (głos + obraz) → wymagany OS (Linux/Android) → kamery HD i przetwarzanie wideo

Tabela 4. Typowe parametry MCU w zastosowaniach TinyML. Dane: Heydari S., Mahmoud Q.H., Sensors 2025, 25(10), 3191

Parametr	Typowy zakres	Przykładowa platforma
Częstotliwość taktowania	100...480 MHz	ARM Cortex-M4/M7, M85
Pamięć Flash (model + kod)	do 1 MB	STM32, Nordic nRF, RA8
Pamięć SRAM (dane robocze)	128...512 kB	Cortex-M4 typowo 256 kB
Pobór mocy (inference)	25...300 mW	zależnie od taktowania
Latencja inferencji	0,18...300 ms	zależy od modelu i kompresji
Akcelerator zmiennoprzecinkowy (FPU)	opcjonalny	Cortex-M4F: sprzętowy FPU
Wbudowany NPU/SIMD	nowe generacje	Helium (Cortex-M85, Armv8.1-M)

Co to jest CoreMark i dlaczego ważny?

CoreMark to benchmark przemysłowy (EEMBC) mierzący wydajność rdzenia CPU przy typowych operacjach embedded: listy, macierze, operacje bitowe, CRC.

Wyniki orientacyjne (przy 1 MHz taktowania):

- Cortex-M0+ ~ 0,9 CoreMark/MHz
- Cortex-M4F ~ 1,9 CoreMark/MHz
- Cortex-M7 ~ 5,0 CoreMark/MHz
- Cortex-M85 ~ 6,3 CoreMark/MHz (Helium)

Wyższa wartość oznacza więcej operacji na jednostkę mocy obliczeniowej, co przekłada się bezpośrednio na możliwości modeli ML możliwych do uruchomienia.

co przekłada się na skokowy wzrost wydajności operacji DSP i ML. Mikrokontrolery zbudowane na tej architekturze osiągają ponad 3000 punktów CoreMark przy taktowaniu 480 MHz – wynik niewyobrażalny dla Cortex-M4 działającego przy tej samej częstotliwości.

Akceleratorzy AI: co wnosi dedykowany sprzęt?

Rdzenie ARM ogólnego przeznaczenia są wszechstronne, ale inferencja sieci neuronowych to w 80...90% mnożenie macierzy. Dedykowane akceleratorzy AI – Neural Processing Units (NPU) lub Google Edge TPU jako znany przykład klasy ASIC – są projektowane wyłącznie pod ten typ obliczeń. Efekty są mierzalne: jak dokumentuje Heydari i Mahmoud, wyspecjalizowane akceleratorzy osiągają wydajność do 20 razy wyższą niż MCU ogólnego przeznaczenia przy zachowaniu poboru mocy w granicach akceptowalnych dla zasilania bateryjnego.

Różnica architektury jest fundamentalna. Klasyczny procesor jest zoptymalizowany pod kątem sekwencyjnego wykonywania kodu z gałęziami warunkowymi i różnorodnymi typami operacji. NPU to w istocie macierz jednostek MAC (Multiply-Accumulate) zdolna do masowego równoległego przetwarzania tensorów. Przy inferencji modelu

ResNet-50 na MCU ogólnego przeznaczenia czas wynosi sekundy; na dedykowanym NPU – milisekundy, przy ułamku zużywanej energii.

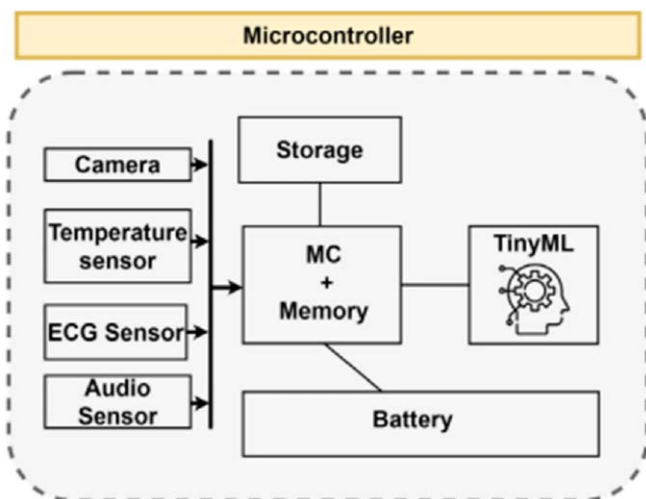
Segment FPGA zajmuje interesującą niszę między elastycznością MCU a wydajnością ASIC. Xilinx Versal, będący przykładem platformy klasy FPGA z wbudowanymi rdzeniami AI Engine, pozwala na implementację modeli ML z możliwością późniejszej rekonfiguracji – bez konieczności zmiany układu. To szczególnie cenne w fazach R&D, kiedy architektura modelu nadal ewoluuje, oraz w niskoseryjnych wdrożeniach przemysłowych, gdzie opłacalność produkcji własnego ASIC jest wątpliwa.

Co-design: dlaczego sprzęt i model muszą rosnąć razem

Jednym z najkosztowniejszych błędów w projektach embedded AI jest myślenie sekwencyjne: najpierw projektuję hardware, potem wybieram model, potem optymalizuję. Podejście to prowadzi do odkrycia – często tuż przed finalną weryfikacją projektu – że wybrany MCU nie ma dość pamięci dla modelu, który okazał się niezbędny do osiągnięcia wymaganej dokładności.

Co-design sprzętu i algorytmu to paradygmat, w którym oba wymiary są optymalizowane równolegle od pierwszego dnia projektu. W praktyce oznacza to, że specyfikacja wymagań modelu ML – rozmiar, latencja, dokładność, zużycie energii – wchodzi do założeń projektanta sprzętu razem ze specyfikacją elektryczną i mechaniczną. Jak ujmują to badania z przeglądu Edge AI (arXiv 2510.01439): hardware, oprogramowanie i warstwa aplikacyjna mają krytyczne współzależności, których nie można analizować w izolacji.

Dowód na to, że co-design jest możliwy nawet w ekstremalnych ograniczeniach zasobowych, dają badania nad on-device training – nie tylko inferencją, ale pełnym uczeniem maszynowym na mikrokontrolerze. Adhikary i współpracownicy w swojej pracy TinyWolf z 2024 roku wykazali, że modyfikacja algorytmu Grey Wolf Optimizer pod kątem zarządzania pamięcią (podział wagi między Flash a SRAM z stronicowaniem) pozwoliła na trening modelu głębszej sieci neuronowej na mikrokontrolerze z zaledwie



Rysunek 2. Architektura sprzętowa TinyML. Wybór komponentów zależy od wymagań zadania (klasy modelu ML) i dostępnych zasobów sprzętowych (pamięć, moc obliczeniowa, budżet energetyczny)

Co-design w praktyce – lista kontrolna

Na etapie specyfikacji projektu (przed wyborem MCU) ustal:

1. Maksymalny rozmiar modelu po kwantyzacji (Flash)
Reguła: zarezerwuj $\geq 50\%$ Flash na model + wagi
2. Wymagana pamięć robocza podczas inferencji (SRAM)
Reguła: aktywacje warstwy = szerokość $\times$ wysokość $\times$ kanały $\times$ precyzja
3. Wymagana latencja pojedynczej inferencji
Reguła: jeśli < 5 ms $\rightarrow$ rozważ MCU z NPU lub SIMD
4. Budżet energetyczny na jeden cykl inferencji (mJ)
Reguła: moc [mW] $\times$ czas inferencji [ms] = energia [μ J]
5. Częstotliwość aktualizacji modelu w polu (OTA?)
Reguła: jeśli TAK $\rightarrow$ zestaw rezerwę Flash $\geq 2 \times$ rozmiar modelu

Tabela 5. Techniki kompresji modeli TinyML i ich efekty. Dane: Heydari i Mahmoud, Sensors 2025

Technika	Mechanizm działania	Typowy efekt	Koszt dokładności
Kwantyzacja (Quantization)	Redukcja precyzji wag: FP32 → INT8 lub mniej	Kompresja 4×, przyspieszenie 2...4×	Nieznaczny (< 1 pp)
Przycinanie sieci (Pruning)	Usuwanie mało istotnych połączeń i neuronów	Do 13× redukcja sieci	Umiarkowany, eliminowalny przez re-trening
Destylacja wiedzy (Knowledge Distillation)	Trenowanie małej sieci na wyjściach dużej	Model 10...50× mniejszy	Zależy od jakości modelu-nauczyciela
Kombinacja technik	Kwantyzacja + pruning równocześnie	Do 49× łączna kompresja	Wymaga starannej optymalizacji

256 kB RAM, oszczędzając 71% pamięci w porównaniu z klasycznymi optymalizatorami gradientowymi. Zużycie energii podczas treningu dla najmniejszej konfiguracji (10 wółfów) mieściło się w zakresie 0,035...0,157 mWh.

To przykład, który zmienia perspektywę: granica między urządzeniem zdolnym wyłącznie do inferencji a urządzeniem zdolnym do nauki w terenie jest coraz mniej ostra. On-device learning to nie odległa przyszłość – to kategoria, która zaczyna wchodzić do produkcyjnych wdrożeń.

Trzy ery kompresji modeli: jak zmieścić sieć w kilobajtach

Nawet najlepszy sprzęt nie wystarczy, jeśli model jest za duży. Kompresja modeli to dyscyplina, która rozwinęła się równolegle z hardware – i bez której era TinyML nie byłaby możliwa. Trzy fundamentalne techniki są dziś standardem w warsztacie każdego inżyniera wdrażającego AI w systemach embedded.

Kwantyzacja jest zwykle pierwszym krokiem, bo jej stosunek zysku do ryzyka jest najkorzystniejszy. Zamiana 32-bitowych liczb zmiennoprzecinkowych na 8-bitowe całkowite (INT8) zmniejsza rozmiar modelu czterokrotnie i przyspiesza inferencję, ponieważ większość nowoczesnych MCU i NPU ma zoptymalizowane ścieżki dla operacji całkowitoliczbowych. Utrata dokładności jest zazwyczaj poniżej jednego punktu procentowego – nieistotna w większości zastosowań przemysłowych.

Przycinanie sieci pozwala zredukować złożoność modelu przez usuwanie połączeń, których wagi są bliskie zeru. Badania cytowane przez Heydari i Mahmoud wykazały, że sieci mogą być zredukowane nawet 13-krotnie bez istotnej utraty dokładności – pod warunkiem iteracyjnego procesu: przytnij → dotrenuj → oceń → przytnij ponownie. Kluczowe jest tu słowo iteracyjnie – jednorazowe agresywne pruning bez re-treningu prowadzi do katastrofalnych strat w jakości modelu.

Optymalizacja pipeline'u kompresji przebiega zwykle w sekwencji: najpierw kwantyzacja (najprostszy zysk), następnie pruning jeśli kwantyzacja nie wystarczyła, wreszcie re-trening po każdym etapie. Praktyczny przegląd z 2025 roku opublikowany w MDPI Electronics potwierdza tę sekwencję jako uznaną metodologię wdrożeniową:

„Quantization is applied first because it is a straightforward method that significantly reduces model size and accelerates inference. If quantization alone is insufficient, the process continues with pruning.”

Podsumowanie

Sprzęt dla Edge AI ewoluował przez trzy wyraźne ery – od wczesnych prób kompresji modeli chmurowych, przez eksplozję mobile AI, po obecną erę TinyML z mikrowatową inferencją na mikrokontrolerach. Dziś projektant ma do wyboru cztery klasy sprzętu o różnych profilach: MCU (niski koszt, niski pobór mocy, ograniczone zasoby), FPGA (rekonfigurowalność, średni koszt), ASIC/NPU (maksymalna wydajność energetyczna, duże serie) i SoC (złożone aplikacje multimedialne, wbudowany OS).

Typowy MCU dla TinyML pracuje z taktowaniem 100...480 MHz, pamięcią SRAM 128...512 kB i poborem mocy 25...300 mW. Dedykowane akceleratory AI oferują do 20-krotną przewagę wydajnościową. Techniki kompresji modeli – kwantyzacja, pruning, destylacja – pozwalają na redukcję rozmiarów do 49× przy zachowaniu użytecznej dokładności.

Kluczowa lekcja tego rozdziału: sprzęt i model ML muszą być projektowane razem, nie sekwencyjnie. Co-design, w którym wymagania modelu kształtują specyfikację hardware już na etapie założeń projektowych, jest warunkiem koniecznym efektywnego embedded AI.

JS

Źródła:

- [1] Huang X., Wang H., Qin S., Tang S.-K. Embedded Artificial Intelligence: A Comprehensive Literature Review. Electronics 2025, 14(17), 3468. doi: 10.3390/electronics14173468 – trzy ery EAI (rozdz. 1), platformy sprzętowe (rozdz. 3).
- [2] Heydari S., Mahmoud Q.H. Tiny Machine Learning and On-Device Inference: A Survey. Sensors 2025, 25(10), 3191. doi: 10.3390/s25103191 – parametry MCU (rozdz. 2.3), akceleratory AI (rozdz. 2.3), techniki kompresji (rozdz. 2.2).
- [3] Adhikary S., Dutta S., Dwivedi A.D. TinyWolf – Efficient On-Device TinyML Training for IoT Using Enhanced Grey Wolf Optimization. Internet of Things 2024, 28, 101365. doi: 10.1016/j.iot.2024.101365 – on-device training na 256 kB RAM (rozdz. 2.5).
- [4] Abou Ali M., Dornaika F. Edge Artificial Intelligence: A Systematic Review of Evolution, Taxonomic Frameworks, and Future Horizons. arXiv 2510.01439, 2025 – taksonomia wielowymiarowa Edge AI, współzależności hardware-software-aplikacja.
- [5] García-Hernández A. et al. Edge AI in Practice: A Survey and Deployment Framework for Neural Networks on Embedded Systems. Electronics 2025, 14(24), 4877. doi: 10.3390/electronics14244877 – metodologia wdrożeniowa, pipeline kompresji.

Synteza dźwięku, część 8

Filtry sterowane napięciem (2)

W poprzednim odcinku omówiony został dolnoprzepustowy filtr sterowany napięciem o całkiem przyzwoitych parametrach. W tej części przyjrzymy się innej konstrukcji filtra. Nie ma bowiem jednego, idealnego rozwiązania, a każda konstrukcja może brzmieć nieco inaczej wnosząc do syntezy swój własny koloryt.

Układ omawiany w tym artykule jest filtrem uniwersalnym zapewniającym trzy różne filtry sterowane wspólnym sygnałem: filtr dolnoprzepustowy (Low Pass – LP), górnoprzepustowy (High Pass – HP) oraz pasmowoprzepustowy (Band Pass – BP). Współczynnik dobroci (Q) regulowany jest potencjometrem, podczas gdy częstotliwość odcięcia jest kontrolowana sygnałem CV. Filtr ten jest filtrem drugiego rzędu, co oznacza nachylenie zboczy 12 dB na oktawę. Sam układ jest relatywnie prosty, dlatego autor tego projektu sugeruje użycie przynajmniej dwóch takich filtrów. Warto dodać też, iż podkreślając współczynnik dobroci filtr zacznie oscylować w okolicy częstotliwości odcięcia generując bardziej interesujące brzmienia.

State Variable Filter – zasada działania

Rysunek 1 przedstawia uproszczony schemat układu State Variable Filter (SVF). Układ ten zawiera trzy wzmacniacze operacyjne połączone ze sobą. Pierwszy stopień zbudowany wokół układu U1 to klasyczny wzmacniacz sumujący dający na wyjściu sygnał filtra

górnoprzepustowego. Sygnał ten trafia na drugi stopień, zbudowany wokół układu U2, który to stopień pracuje jako układ całkujący. Wyjście tego stopnia daje sygnał pasmowoprzepustowy. Sygnał ten trafia na wejście nieodwracające pierwszego stopnia, ale też i na wejście stopnia trzeciego, który jest identyczny ze stopniem drugim, i używając operacji całkowania przekształca ten sygnał na sygnał filtra dolnoprzepustowego. Sygnał ten jest dodawany do sygnału wejściowego i trafia na wejście odwracające układu U1, i zostaje odjęty od sygnału filtra pasmowoprzepustowego, co w efekcie daje sygnał górnoprzepustowy. Zakładając, iż $R_{f1}=R_{f2}$, $R_1=R_2$, oraz $C_1=C_2$, częstotliwość odcięcia wynosi:

$$f_0 = \frac{1}{2\pi R_{f1} C_1}$$

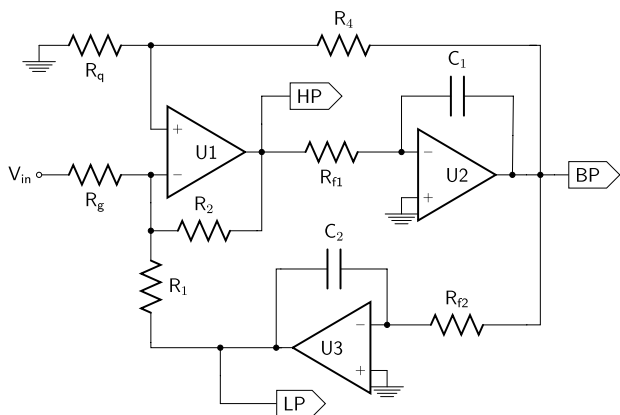
a współczynnik dobroci wynosi:

$$Q = \left(1 + \frac{R_4}{R_q}\right) \left(\frac{1}{2 + \frac{R_1}{R_g}}\right)$$

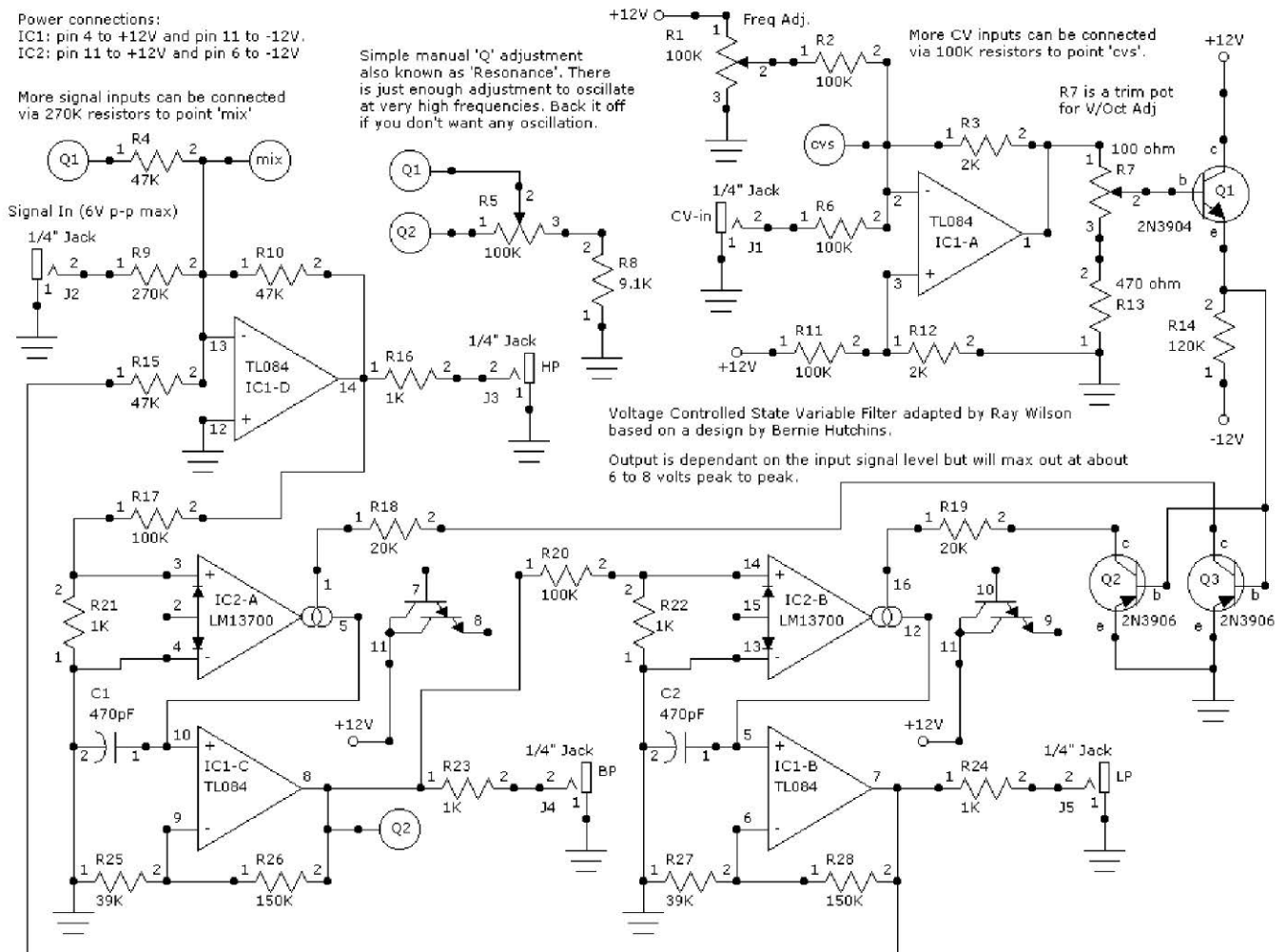
Z tych wzorów wynika, iż częstotliwość odcięcia i współczynnik dobroci są od siebie niezależne. Wzmocnienie dla pasma przenoszenia wyjść dolnoprzepustowego i górnoprzepustowego ustala stosunek wartości R_1 do R_g . Niezależna kontrola nad częstotliwością odcięcia, współczynnikiem dobroci i wzmocnieniem filtra sprawiają, iż jest to niezwykle użyteczny układ, szczególnie tam, gdzie potrzebna jest specyficzna dobroć – inne filtry aktywne wyższych rzędów nie pozwalają na tak dokładną jej kontrolę. By uczynić ten filtr użytecznym dla syntezy, trzeba go nieco zmodyfikować.

Praktyczny filtr SVF

Rysunek 2 przedstawia realizację filtra SVF projektu Raya Wilsona [1] na bazie wcześniejszego projektu Berniego Hutchinsa. Układ zawiera wszystkie podstawowe elementy filtra uniwersalnego: wzmacniacz sumujący na wejściu i dwa stopnie całkujące, choć na pierwszy rzut oka nie wyglądają one na stopnie całkujące, tylko różniczkujące. Standardowo kondensator we wzmacniaczu całkującym znajduje się w pętli sprzężenia zwrotnego, tu jednak jest umieszczony w szeregu z jednym z wejść wzmacniacza. Jednakże jest to nadal układ całkujący dzięki obecności wzmacniacza transkonduktancyjnego w obwodzie kondensatora.



Rysunek 1. Filtr typu State Variable Filter (SVF) – schemat ogólny.
Źródło: Wikipedia



Rysunek 2. Praktyczny filtr SVF o relatywnie prostej konstrukcji

Sygnał wejściowy przez rezystor R9 trafia na wejście odwracające wzmacniacza TL084 IC1-D, do którego też dociera sygnał z wyjścia filtra dolnoprzepustowego. Do tego sygnału dodawany jest też sygnał z wyjścia filtra pasmowoprzepustowego przez potencjometr R5 rezystor R4. Potencjometr oraz R8 tworzą dzielnik ustalający współczynnik dobroci układu. Sygnał z IC1-D trafia przez dodatkowy rezystor R17 na bocznik R21. Bocznik ten włączony jest między wejścia wzmacniacza transkonduktancyjnego IC2-A, a różnica między nimi określa prąd wyjściowy, który trafia na kondensator C1. Napięcie na tym kondensatorze zależy od jego pojemności, prądu ładowania/rozładowywania (zależnego od różnicy napięć na R21 i prądu sterującego I_{abc}) i prędkości zmian tego prądu. Kondensator realizuje w ten sposób funkcję całkowania. Napięcie na tym kondensatorze jest buforowane przez IC1-C, którego wzmocnienie ustala stosunek R25 do R26. Na wyjściu powstaje sygnał filtra pasmowoprzepustowego, który trafia do wyjścia oraz do kolejnego stopnia całkującego zrealizowanego z użyciem IC2-B, IC1-B i kondensatora C2. Ten układ generuje sygnał filtra dolnoprzepustowego.

Sygnał sterujący CV trafia przez R6 na wejście IC1-A, który też pełni funkcję sumatora oraz bufora. Do tego sygnału

dodawany jest drugi sygnał z potencjometru R1 pozwalającego doregulować częstotliwość filtra. Trymer R7 i rezystor R13 pozwalają wyregulować filtr, by zmiana napięcia na wejściu o 1 V zmieniała częstotliwość odcięcia o jedną oktawę. Sygnał z trymera steruje bazą tranzystora Q1, który jest podłączony w układzie wtórnika emiterowego, i który kontroluje bazy Q2 i Q3 pracujących w układzie wspólnego emitera. Tranzystory te kontrolują prądy I_{abc} obu wzmacniaczy transkonduktancyjnych. Całość zasilana jest napięciem symetrycznym ± 12 V.

Podsumowanie

Filtry SVF są relatywnie prostymi konstrukcjami i oferują dużą swobodę tworzenia brzmień w syntezaźach modularnych. Ray Wilson sugeruje wręcz posiadanie przynajmniej dwóch takich filtrów ze względu na ich prostotę. W następnej części przyjrzymy się jeszcze jednemu, rozbudowanemu filtrowi.

Paweł Kowalczyk

Źródła:

[1] https://musicfromouterspace.com/analogsynth_new/OLDIESBUTGOODIES/VCF/vcfstatevar.html

Oscyloskopy 2026 – od lampy Brauna do 12 bitów na biurku konstruktora, część 2

W pierwszej części repetytorium opublikowanej w poprzednim wydaniu EP omówiliśmy historię rozwoju oscyloskopów (rozdział 1) oraz architekturę współczesnego DSO – tor pionowy, przetwornik ADC i pamięć akwizycji (rozdział 2). Pokazaliśmy schemat blokowy oscyloskopu cyfrowego, omówiliśmy konsekwencje wyboru rozdzielczości pionowej 8 vs 10 vs 12 bitów i wyjaśniliśmy, dlaczego R&S MXO osiąga ponad 4,5 miliona przebiegów na sekundę dzięki dedykowanej architekturze akwizycji.

Druga część, którą prezentujemy teraz, przesuwa akcent z oscyloskopu jako sprzętu na sam pomiar. Rozdział 3 omawia próbkowanie, wyzwalanie i akwizycję w praktyce – dlaczego deklaracja „1 GSa/s w karcie katalogowej” nie zawsze oznacza 1 GSa/s na każdym kanale, jak rozpoznać aliasing na ekranie, kiedy włączyć tryb hi-res a kiedy peak detect, oraz jak wyzwalanie zaawansowane (pulse width, runt, time-out, B-trigger) potrafi w ciągu sekund wyłapać błąd, który dla klasycznych triggerów zboczowych byłby praktycznie niewidoczny. Rozdział 4 schodzi na poziom sondy i front-endu analogowego: pokazujemy konstrukcję sond pasywnych 1:10, kiedy warto sięgnąć po sondę aktywną lub różnicową, dlaczego pętla masy o długości 15 cm potrafi przekłamać pomiar 100 MHz bardziej niż 20% niższe pasmo samego oscyloskopu, i jak interpretować kategorie bezpieczeństwa CAT II/III/IV.

W trzeciej, ostatniej części cyklu, planowanej do publikacji w kolejnym wydaniu EP, przyjrzymy się trendowi,

który w ciągu ostatnich pięciu lat najmocniej zmienił ofertę oscyloskopów: wejściu 12-bitowych ADC i wbudowanych kanałów logicznych (MSO) do klasy budżetowej, czyli poniżej 10 000 zł (rozdział 5). Zamknijemy cykl rozdziałem 6, w którym zestawimy reprezentatywne modele 2026 r. ze szczegółową tabelą porównawczą i praktycznymi wskazówkami doboru oscyloskopu do konkretnego zastosowania w pracowni konstruktora.

Materiał podobnie jak w części 1 opieramy na klasycznych dokumentach producenckich – Keysight „Basic Oscilloscope Fundamentals” (AN 5989-8064EN), Tektronix „XYZs of Oscilloscopes” (03W-8605), R&S „Oscilloscope Fundamentals” – uzupełnionych o materiały szczegółowe dotyczące tematyki rozdziałów 3 i 4: Keysight „Evaluating Oscilloscope Sample Rate vs. Sampling Fidelity” (AN 5989-5732), dokumenty Keysight i Tektronix o trybie segmentowym (5989-7833EN, 5989-4932EN, 55W-12112, 55W-61299), R&S „Oscilloscope Basics” z rozdziałem 7 poświęconym wyzwalaniu, Tektronix „ABCs of Probes” (60W-6053-15) i Keysight „Eight Hints for Better Scope Probing” (AN 5989-7894EN).

Rozdział 3. Próbkowanie, wyzwalanie i akwizycja w praktyce konstruktora

W rozdziale 2 omówiliśmy schemat blokowy DSO i sześć bloków, przez które przepływa sygnał. W tym rozdziale skupimy się na trzech najbardziej krytycznych etapach tej drogi: na próbkowaniu (twierdzenie Nyquista, aliasing, fidelity), na wyzwalaniu (od prostego edge triggera do warunków sekwencyjnych) i na trybach akwizycji (sample, peak detect, hi-res, average, envelope, segmented memory). Cel jest praktyczny: pokazać,

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 – bieżące wydanie)

- **Rozdział 3.** Próbkowanie, wyzwalanie i akwizycja – tryby hi-res i pamięć segmentowa, wyzwalanie zaawansowane.
- **Rozdział 4.** Sondy i integralność sygnału – sondy pasywne i aktywne, pętla masy, kategorie bezpieczeństwa.

Część 3 (następne wydanie EP 07-08)

- **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.

dla czegoś deklaracja „1 GSa/s w karcie katalogowej” nie zawsze oznacza 1 GSa/s na każdym kanale, jak rozpoznać aliasing na ekranie, kiedy włączyć tryb hi-res, a kiedy peak detect, oraz jak wyzwalanie zaawansowane potrafi w ciągu sekund wyłapać błąd, który dla klasycznym triggera zboczowego byłby praktycznie niewidoczny.

Materiał opieramy na tych samych tutorialach producentów co rozdział 2 – Keysight „Basic Oscilloscope Fundamentals” (AN 5989-8064EN), Tektronix „XYZs of Oscilloscopes” (03W-8605) i R&S „Oscilloscope Fundamentals” – uzupełnionych o materiały szczegółowe: Keysight „Evaluating Oscilloscope Sample Rate vs. Sampling Fidelity” (AN 5989-5732) do sekcji 3.1...3.2, dokumenty Keysight 5989-7833EN i 5989-4932EN oraz Tektronix 55W-12112 i 55W-61299 do sekcji o trybie segmentowym (3.6), oraz R&S „Oscilloscope Basics” (rozdział 7 dedykowany wyzwalaniu) do sekcji 3.4...3.5.

3.1. Twierdzenie Nyquista – co właściwie znaczy „dwa razy pasmo”

Klasyczne twierdzenie Shannona-Nyquista mówi, że aby sygnał o ograniczonym widmie można było zrekonstruować bezstratnie z próbek, częstotliwość próbkowania musi być co najmniej dwukrotnie wyższa od najwyższej składowej częstotliwościowej sygnału. Dla sygnału o paśmie 100 MHz minimalna częstotliwość próbkowania to 200 MSA/s. Rekonstrukcja w tym przypadku jest matematyczna – z próbek punktowych odtwarza się pierwotną falę za pomocą interpolacji sinc.

W oscyloskopach inżynierskich kryterium 2× obowiązuje teoretycznie, ale praktyka jest inna. Producenci rekomendują 2,5...5× pasmo analogowe oscyloskopu (Keysight w AN 5989-5732 omawia to szczegółowo). Powód jest prosty: interpolacja sinc jest dokładna tylko dla nieskończonej długości rekordów próbek i przy idealnych filtrach antyaliasingowych. W rzeczywistych oscyloskopach akwizycja jest ograniczona długością pamięci, a filtr antyaliasingowy nie jest idealny – ma skończone nachylenie zbocza. Dlatego dla sygnału cyfrowego o paśmie 500 MHz Keysight zaleca oscyloskop pracujący przy 5 GSA/s (10×), a nie 1 GSA/s (2×). Niższe próbkowanie da się obronić matematycznie, ale na ekranie zobaczymy wówczas zarys obwiedni zamiast kształtu zboczy.

Trzeba też pamiętać o tym, że „pasmo” w karcie katalogowej oscyloskopu to częstotliwość graniczna toru analogowego (–3 dB), nie limit twardy. Sygnał o częstotliwości 1,5× pasma jest wciąż widziany przez oscyloskop, tylko z tłumieniem rzędu 6...10 dB i z lekko zniekształconym kształtem zboczy. W praktyce nie wolno więc rozważać pasma jako filtra cliff-cut – jest to miękka granica, za którą zaczynają się błędy amplitudy.

3.2. Aliasing – kiedy zobaczymy go na ekranie

Aliasing to zjawisko fundamentalne: kiedy częstotliwość próbkowania spadnie poniżej dwukrotności częstotliwości sygnału, próbki układają się na falowej obwiedni o niższej częstotliwości – ale oscyloskop pokazuje tę obwiednię jako „prawdziwy” sygnał, nie ostrzega o aliasingu. Wzór, opisujący zjawisko jest prosty:

$$f_{alias} = |f_{sygnał} - k \cdot f_{próbk}| \quad (\text{dla całkowitego } k)$$

W oscyloskopach najczęściej widzimy alias pierwszego rzędu, czyli dla $k = 1$.

Przykład praktyczny pokazuje **rysunek 10**. Ten sam sygnał sinusoidalny 100 MHz próbkujemy najpierw z częstotliwością 1 GSA/s (panel A, 10× pasmo sygnału – kryterium Nyquista spełnione z dużym zapasem), a następnie z 110 MSA/s (panel B, poniżej kryterium 2×). W pierwszym przypadku próbki gęsto pokrywają sygnał i cyfrowa rekonstrukcja odtwarza 100 MHz wiernie. W drugim przypadku próbki łapią sygnał w niesynchroniczne miejsca, a oscyloskop wyświetla wolniejszą sinusoidę o częstotliwości $|100 \text{ MHz} - 110 \text{ MHz}| = 10 \text{ MHz}$. Sygnał 100 MHz jest niewidoczny, a alias 10 MHz wygląda jak prawdziwy przebieg.

Najlepszy mechanizm obronny przed aliasingiem w praktyce to: (1) wybierać sample rate co najmniej 2,5× powyżej pasma analogowego oscyloskopu (a nie spodziewanego pasma sygnału – realny sygnał może mieć szybkie przebiegi przejściowe, których istnienia nie podejrzewamy), (2) zostawić oscyloskop w trybie auto sample rate, nie zmuszać go do niskiej szybkości próbkowania ręcznie, (3) jeżeli widzimy na ekranie podejrzenie wolny ślad – zmienić timebase i zobaczyć, czy „wolna” częstotliwość zmienia się przy zmianie sample rate. Prawdziwy sygnał trzyma swoją częstotliwość przy zmianie podstawy czasu, alias się zmienia.

3.3. Tryby akwizycji – sample, peak detect, hi-res, average, envelope

Każdy oscyloskop klasy inżynierskiej oferuje co najmniej pięć trybów akwizycji. Wybór trybu decyduje, co dokładnie trafia z ADC do pamięci akwizycji. Większość użytkowników zostawia tryb Sample (domyślny) i nigdy go nie zmienia – to błąd, bo każdy z czterech pozostałych trybów dramatycznie zmienia to, co widać na ekranie.

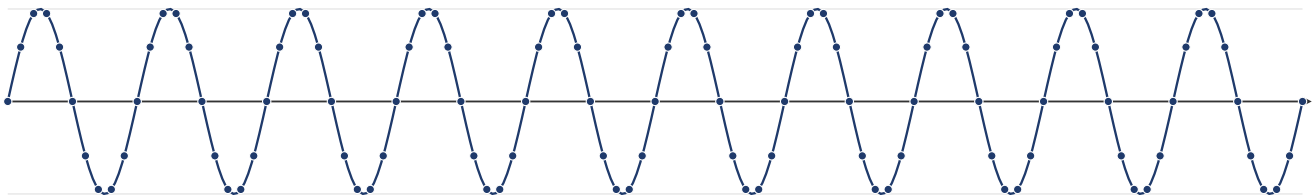
Tryb Sample (próbkowanie surowe) zapisuje każdą próbkę z ADC bez modyfikacji. To jest baseline – bez post-processingu, bez uśredniania, bez decymacji. Wszystkie pozostałe tryby pracują na strumieniu próbek wychodzących z trybu Sample i go modyfikują.

Peak detect to tryb, w którym oscyloskop przy zmniejszonej podstawie czasu (czyli przy długim oknie obserwacji) zamiast decymacji zapisuje dwie próbki na każdy

Aliasing: dlaczego próbkowanie poniżej kryterium Nyquista zniekształca sygnał

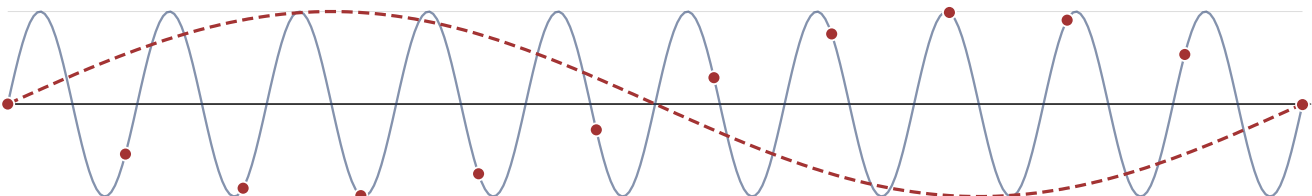
Ten sam sygnał 100 MHz próbkowany z dwiema różnymi częstotliwościami

A. Sample rate = 1 GSa/s ($10\times$ pasmo sygnału) — sygnał odtworzony poprawnie



Próbki gęsto pokrywają sygnał → cyfrowa rekonstrukcja odtwarza 100 MHz wiernie.

B. Sample rate = 110 MSa/s (poniżej kryterium Nyquista) — alias 10 MHz zamiast 100 MHz



Próbki łapią sygnał w niesynchroniczne miejsca → oscyloskop rekonstruuje wolniejszą sinusoidę (alias) o częstotliwości $|f_{\text{sygnał}} - f_{\text{próbk.}}| = |100 \text{ MHz} - 110 \text{ MHz}| = 10 \text{ MHz}$. Sygnał 100 MHz jest niewidoczny.

Legenda:

— sygnał wejściowy (100 MHz)

• próbki przy $f_s > 2 \cdot f_{\text{sygnał}}$ (poprawnie)

--- alias rekonstruowany przez oscyloskop (10 MHz)

• próbki przy $f_s < 2 \cdot f_{\text{sygnał}}$ (aliasing)

Kryterium Nyquista: do bezstratnej rekonstrukcji sygnału próbkowanie musi być co najmniej dwukrotnie wyższe niż najwyższa składowa częstotliwościowa sygnału.

Źródło: opracowanie graficzne EP, wzorowane na: Keysight Technologies, „Evaluating Oscilloscope Sample Rate vs. Sampling Fidelity”, Application Note 5989-5732, ilustracje aliasingu w funkcji częstotliwości próbkowania; Tektronix „XYZs of Oscilloscopes” (03W-8605), rozdz. Sampling.

Rysunek 10. Aliasing przy niedostatecznej częstotliwości próbkowania. Panel A: $f_s = 1 \text{ GSa/s}$ ($10\times$ pasmo sygnału) – sygnał 100 MHz odtworzony poprawnie. Panel B: $f_s = 110 \text{ MSa/s}$ (poniżej kryterium Nyquista, które wymaga $\geq 200 \text{ MSa/s}$) – oscyloskop rekonstruuje alias 10 MHz zamiast rzeczywistego sygnału 100 MHz. Źródło: opracowanie graficzne EP, wzorowane na: Keysight Technologies, „Evaluating Oscilloscope Sample Rate vs. Sampling Fidelity”, Application Note 5989-5732, ilustracje aliasingu w funkcji częstotliwości próbkowania; Tektronix „XYZs of Oscilloscopes” (03W-8605), rozdz. Sampling

piksel ekranu: minimum i maksimum z całego okna pikselowego. Po co? Wyobraźmy sobie, że oglądamy zasilanie SMPS na podstawie czasu 200 ms/dz, czyli 2 sekundy na ekran. Przy 1 GSa/s mielibyśmy w pamięci 2 mld próbek, ale głębokość pamięci to typowo 10...100 Mpts. Oscyloskop musi zatem decymować próbki lub zmniejszyć sample rate. W trybie Sample krótki impuls glitch o długości 50 ns (a więc 50 próbek z ADC) najprawdopodobniej zostanie zgubiony w decymacji – wypadnie między „klatkami” wyświetlanymi na ekranie. W trybie peak detect oscyloskop monitoruje cały strumień próbek z ADC i zapisuje minimum i maksimum – glitch pojawi się na ekranie jako wąski pionowy ślad, bo jego maksimum było wyższe niż otaczające próbki. Peak detect to tryb pierwszego wyboru dla obserwacji „czy coś się dzieje” na długich oknach czasowych.

Hi-res to tryb opisany w rozdziale 2.3 – oversampling z uśrednianiem, który pozwala uzyskać 12-bit rozdzielczość z natywnie 8-bitowego ADC kosztem redukcji

pasma. Wpisując to w karty katalogowe producenci często stosują skrót: „12-bit w trybie hi-res” oznacza właśnie ten mechanizm. Tryb hi-res jest sensowny dla powolnych sygnałów (zasilacze, czujniki, sygnały audio), niesensowny dla szybkich przebiegów przejściowych.

Average to tryb, w którym oscyloskop nakłada na siebie N kolejnych akwizycji i wyświetla średnią. Sensowny tylko dla sygnałów powtarzalnych (zsynchronizowanych z triggerem). Szum biały redukuje się o $\sqrt{N}$ – uśrednienie 16 akwizycji daje $4\times$ lepszy SNR, 256 akwizycji daje $16\times$ lepszy SNR. Tryb average wyłącza obserwację zdarzeń jednorazowych, bo to, co nie wystąpi w każdej akwizycji, jest tłumione. Klasyczne zastosowanie: pomiar bardzo małego sygnału RF na tle szumu termicznego przedwzmacniacza albo charakterystyka jitteru zegara.

Envelope to tryb pokrewny peak detectowi, ale działający na poziomie akwizycji, nie pikseli ekranu. Oscyloskop wykonuje N akwizycji i zapisuje obwiednię minimum/maksimum z wszystkich. Pokazuje to „zakres”, w którym sygnał

przebiega – przydatne do mierzenia jitteru (jeżeli zbocze migocze w pewnym oknie czasowym, w trybie envelope zobaczymy pasek o szerokości tego okna), do oceny stabilności sygnału, albo do wyłapania pojedynczych anomalii na tle wielu akwizycji okresu sygnału powtarzalnego.

Praktyczna zasada wyboru trybu: Sample do podstawowych pomiarów, peak detect przy długich oknach czasowych (slow timebase) jeżeli szukamy glitchów, hi-res do powolnych sygnałów wymagających rozdzielczości pionowej, average do redukcji szumu w sygnałach powtarzalnych, envelope do oceny stabilności sygnału w czasie.

3.4. Wyzwalanie podstawowe – źródła, sprzężenia, holdoff

Trigger jest mechanizmem, który decyduje, w którym momencie oscyloskop zatrzyma obraz na ekranie. W każdym oscyloskopie inżynierskim trigger ma trzy parametry kluczowe (klasyfikacja z R&S „Oscilloscope Basics”, rozdz. 7): minimum detectable glitch (najmniejszy impuls, który wyzwoli trigger – typowo 1...10 ns), threshold accuracy (dokładność poziomu progu, typowo $\pm$ kilkanaście mV), oraz trigger jitter (rozrzut czasowy momentu wyzwolenia, typowo kilka ps RMS). Dobry analogowy trigger daje jitter ok. 10 ps, cyfrowy trigger w R&S RTO potrafi zejść poniżej 1 ps RMS.

Najbardziej podstawowy trigger to edge trigger – wyzwalanie zboczem. Konfiguracja składa się z czterech ustawień: źródło (channel 1, channel 2, External, AC line), zbocze (rising, falling, either), poziom (wartość napięcia, na której zatrzaśniemy obraz) i sprzężenie. Dla 90% pomiarów edge trigger wystarcza – nie eksperymentujmy z fancy trygerami, dopóki edge nie zawodzi.

Sprzężenie wyzwalania (trigger coupling) jest często źle rozumiane. Dostępne opcje to: DC (sygnał wyzwalania jest taki sam jak sygnał wyświetlany), AC (przed komparatorem trigger jest filtr górnoprzepustowy ~50 Hz; eliminuje to składową stałą i bardzo wolne dryf), HF reject (filtr dolnoprzepustowy ok. 50 kHz – odcina szum w. cz., dla bardzo wolnych sygnałów z zakłóceniami RF), LF reject (filtr górnoprzepustowy ok. 50 kHz – odcina szum 50 Hz i wolne tętnienia zasilania), oraz noise reject (rozszerza histerezę komparatora; trigger nie wyzwoli się przy małych oscylacjach wokół poziomu progu). Klasyczne błędne ustawienie: chcemy wyzwalać na zboczu zegara mikrokontrolera, sygnał ma DC offset 1,5 V i amplitudę 3,3 V, ustawiamy trigger w trybie AC – wtedy DC offset jest filtrowany, ale szybkie zbocze przechodzi przez filtr górnoprzepustowy bez problemu. Trigger się wyzwala stabilnie. Druga klasyczna sytuacja: trigger się wyzwala chaotycznie na każdym szumie próbnika – włączamy noise reject i komparator wymaga sygnału, który przejdzie przez całą histerezę, więc szum nie wyzwala fałszywie.

Pułapka kart katalogowych – „1 GSa/s” na ile kanałów?

W większości oscyloskopów inżynierskich klasy do 1 GHz pasma deklarowany sample rate w karcie katalogowej (np. 5 GSa/s) jest osiągnięty tylko gdy aktywny jest jeden kanał. Włączenie dwóch kanałów dzieli sample rate na pół (2,5 GSa/s na kanał), czterech kanałów – na cztery (1,25 GSa/s na kanał). Mechanizmem jest interleaving ADC: oscyloskop ma fizycznie cztery przetworniki 1,25 GSa/s, które dla jednokanałowego pomiaru łączą w jeden wirtualny ADC 5 GSa/s. Przy aktywnych wszystkich czterech kanałach każdy z nich pracuje na swoim własnym ADC.

Konsekwencja praktyczna: przy pomiarze trzech sygnałów cyfrowych klasy 200 MHz w trzech kanałach oscyloskopu 1 GSa/s nie mamy 1 GSa/s na każdym z nich – mamy efektywnie 250 MSa/s (kryterium Nyquista wciąż jest spełnione dla 200 MHz, ale ledwo). Karta katalogowa zawsze podaje obie wartości: „5 GSa/s (1 kanał)/1,25 GSa/s (4 kanały)” albo „interleaved sample rate” – warto czytać uważnie i dobrać oscyloskop do najgorszego przypadku swojej aplikacji, nie do najlepszego.

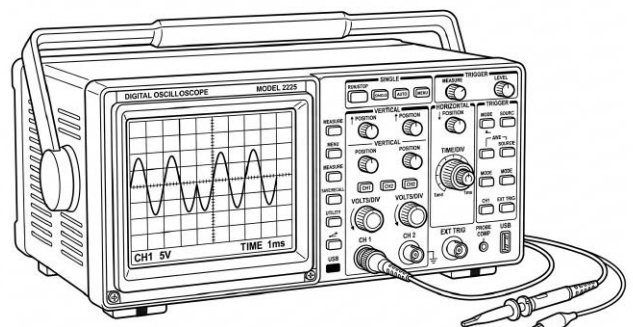
Wyjątkiem są oscyloskopy high-end z dedykowanym przetwornikiem ADC na każdy kanał (R&S MXO, Keysight Infiniium S/Z, Tek MSO 6) – tam każdy kanał ma własny ADC o pełnej szybkości i interleaving nie obowiązuje. Ale to klasa cenowa od 30 000 zł wzwyż, nie na typowe biurko konstruktora.

Holdoff to czas po wyzwoleniu, w którym oscyloskop ignoruje kolejne potencjalne wyzwolenia. Klasyczny przypadek: sygnał składa się z paczek (bursts) impulsów oddzielonych przerwami. Bez holdoff oscyloskop wyzwala się na każdy impuls w paczce i obraz „pływa”. Włączamy holdoff dłuższy od długości paczki – oscyloskop wyzwala się na początku każdej paczki i ignoruje pozostałe impulsy. Drugi przypadek: sygnał ma podwójne zbocze (np. z odbicia na linii transmisyjnej) – holdoff krótszy niż pełen okres, ale dłuższy niż czas trwania odbicia eliminuje wyzwalanie na odbiciu.

3.5. Wyzwalanie zaawansowane – impulsowe, runt, time-out, sekwencyjne

Edge trigger wystarcza w 90% przypadków, ale tych pozostałych 10% to często najtrudniejsze problemy diagnostyczne, w których edge zawodzi z prostej przyczyny: zbocze, którego szukamy, wygląda tak samo jak wszystkie inne zbocza w sygnale. Potrzebujemy bardziej szczegółowych warunków wyzwalania.

Pulse width trigger (wyzwalanie szerokością impulsu) wyzwala się na impuls o szerokości większej,



mniej lub w zakresie. Klasyczne zastosowanie: szukamy glitchów na linii zegara – wszystkie „prawdziwe” impulsy mają szerokość np. 50 ns, glitch ma 5 ns. Ustawiamy pulse width <20 ns – oscyloskop wyzwała się tylko na glitche. Drugie zastosowanie: szukamy zbyt długich impulsów – w protokole I²C bit „start” ma określoną szerokość, anomalia może być impulsem szerszym.

Runt trigger wyzwała się na impulsy, których amplituda przekracza dolny próg, ale nie osiąga górnego. Krytyczne dla diagnostyki interfejsów cyfrowych – „runt” to impuls, który nie osiągnął pełnego poziomu logicznego (np. nie doszedł do 3 V w logice 3,3 V), więc układ odbiorczy go zinterpretuje raz jako „1”, raz jako „0”. Klasyczne źródło sporadycznych, trudnych do zlokalizowania błędów. Edge trigger nigdy nie wyłapie runt, bo zbocze rosnące i opadające wygląda tak samo jak każde inne. Runt trigger jest dedykowany temu problemowi.

Time-out trigger (wyzwalanie czasem braku zdarzenia) wyzwała się wtedy, gdy sygnał nie miał zbocza przez

zadany czas. Klasyczne zastosowanie: linia heartbeat lub watchdog – sygnał powinien mieć zbocze co 100 ms, jeśli nie miał przez 200 ms to znaczy, że hardware się zawiesił. Ustawiamy time-out 200 ms i widzimy dokładnie ten moment, w którym sygnał zamilkł. Bez time-outa trzeba by oglądać megabajty zapisanych próbek.

B-trigger (wyzwalanie sekwencyjne) wyzwała się dopiero po spełnieniu warunku A i potem warunku B w określonym oknie czasowym (lub po określonej liczbie zdarzeń). Klasyczne zastosowanie: chcemy zobaczyć stan magistrali SPI w 10-tym cyklu zegara po sygnale chip-select – warunek A to opadające zbocze CS, warunek B to dziesiąte zbocze zegara, dopiero wtedy zatrzymujemy obraz. Bez B-triggera trzeba by ręcznie szukać tego momentu w długim rekordzie.

Logic pattern trigger (wyzwalanie wzorcem logicznym) wyzwała się na zdefiniowany wzorec na wielu kanałach jednocześnie – np. wzorec 1010 na czterech bitach adresu, lub jednocześnie wystąpienie wysokiego stanu

Tryb segmentowy: jak pamięć akwizycji wystarcza na minuty obserwacji

Ten sam strumień rzadkich impulsów zapisany klasycznie i w trybie segmentowym

A. Akwizycja klasyczna — pamięć wypełniona ciągłym zapisem (większość = puste odstępy)

Strumień sygnału (rzeczywisty czas):



Pamięć akwizycji (np. 10 Mpts, jeden ciągły zapis):



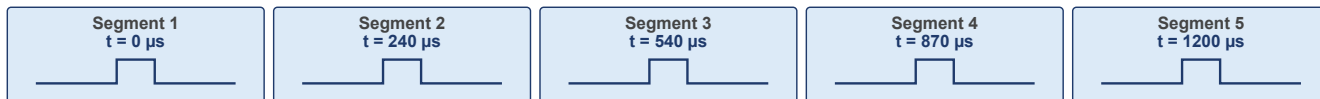
Konsekwencja: 99% pamięci „zajęta” odstępami między impulsami. Aby objąć dłuższy czas, trzeba albo zmniejszyć sample rate (gubimy rozdzielczość czasową impulsu), albo kupić oscyloskop z głębszą pamięcią.

B. Tryb segmentowy — każde wyzwolenie zapisane jako oddzielny segment z etykietą czasową

Ten sam strumień sygnału:



Pamięć akwizycji (taka sama jak A, ale podzielona na 5 segmentów):



Konsekwencja: ta sama pamięć obejmuje 5 wyzwoleń rozciągniętych na 1,2 ms — przy zachowaniu pełnej rozdzielczości czasowej każdego impulsu. Każdy segment ma własną etykietę czasową — różnicę między kolejnymi wyzwoleniami można odczytać z dokładnością wyznaczoną przez zegar akwizycji oscyloskopu.

Zysk z trybu segmentowego:

- Ta sama pojemność pamięci obejmuje setki-tysiące razy dłuższe okno czasowe między pierwszym a ostatnim impulsem.
- Każdy segment zachowuje pełną rozdzielczość czasową (sample rate ADC), więc kształt impulsu pozostaje czytelny.
- Etykiety czasowe segmentów pozwalają zmierzyć odstępy między wyzwoleniami (jitter, częstotliwość burstów, anomalie).

Źródło: opracowanie graficzne EP, wzorowane na: Rohde & Schwarz, „Oscilloscope Fundamentals” (rozdz. „Segmented Memory”); koncepcja segmented memory / FastFrame opisana również w Keysight „InfiniiVision Segmented Memory” oraz Tektronix „FastFrame Application Note”.

Rysunek 11. Porównanie akwizycji klasycznej i trybu segmentowego dla strumienia rzadkich impulsów. W trybie klasycznym (panel A) większość pamięci akwizycji jest zajęta zapisem nieaktywnych odstępów między impulsami. W trybie segmentowym (panel B) ta sama pojemność pamięci dzielona jest na N segmentów aktywowanych przez kolejne wyzwolenia – każdy segment zachowuje pełną rozdzielczość czasową, a etykieta czasowa segmentu pozwala zmierzyć odstęp między wyzwoleniami. Źródło: opracowanie graficzne EP, wzorowane na: Rohde & Schwarz, „Oscilloscope Fundamentals” (rozdział „Segmented Memory”); Keysight Technologies, „Segmented Memory Acquisition for InfiniiVision Series Oscilloscopes”, Data Sheet 5989-7833EN (rewizja 2023); Tektronix, „Using FastFrame Segmented Memory”, Application Note 55W-12112

na trzech sygnałach kontrolnych. W praktyce konstruktora MSO (z kanałami logicznymi) jest to często używane do diagnostyki magistral równoległych.

Wreszcie nowoczesne oscyloskopy MSO oferują dedykowane triggery protokolarne: I²C trigger (start, stop, address, NACK), SPI trigger (chip-select start, slave data pattern), UART/CAN/LIN trigger (frame ID, frame data, frame error). Te triggery działają na poziomie dekodowanego protokołu – zamiast wyzwać na zbocze, wyzwalamy na „pakiet o adresie 0x52 z błędem ACK”. To radykalnie redukuje czas diagnostyki magistral szeregowych.

3.6. Tryb segmentowy – jak pamięć starcza na minuty obserwacji

Tryb segmented memory (FastFrame u Tektroniksa, Segmented Memory u Keysight i R&S) zaadresowany jest do jednego konkretnego problemu: sygnał składa się z krótkich, rzadkich impulsów oddzielonych długimi okresami ciszy. Klasyczna akwizycja w takim przypadku marnuje 99% pamięci na puste odstępy między impulsami. Tryb segmentowy zapisuje wyłącznie krótkie segmenty wokół każdego wyzwolenia, a okresy ciszy są wykluczone z zapisu – pamięć jest dzielona na N segmentów, każdy z własną etykietą czasową (timestamp), które razem obejmują o wiele dłuższe okno obserwacji.

Pokazuje to porównanie z **rysunku 11**. W trybie klasycznym (panel A) pięć rzadkich impulsów oddzielonych długimi przerwami wypełnia całą pamięć akwizycji – ale tylko pięć krótkich fragmentów zawiera rzeczywiście sygnał, reszta to puste tło. W trybie segmentowym (panel B) ta sama pamięć obejmuje te same pięć impulsów, ale zachowuje pełną rozdzielczość czasową każdego z nich, a etykiety czasowe segmentów pozwalają zmierzyć odstępy między wyzwoleniami z dokładnością wyznaczoną przez zegar akwizycji oscyloskopu.

Konkretne liczby: Keysight w karcie katalogowej 5989-7833EN deklaruje dla InfiniiVision do 2000 segmentów z rozdzielczością timestampu 10 ps, a w notcie aplikacyjnej 5989-4932EN podaje przykład z rodziny Infiniium, gdzie wewnętrzny inter-segment time wynosi 2,5 μs dla modeli do 6 GHz, a maksymalna liczba segmentów to 131 072 z opcją 01G. Tektronix w technical brief 55W-61299 dla serii MSO 4/5/6 pokazuje też dodatkową funkcję Summary frame – uśrednioną klatkę zbiorczą na końcu rekordu, co daje od razu uśredniony obraz wszystkich zarejestrowanych zdarzeń. R&S w MXO osiąga podobne liczby segmentów dzięki dedykowanej architekturze akwizycji omawianej w rozdziale 2.5.

Tryb segmentowy nadaje się idealnie do trzech klas zastosowań w warsztacie konstruktora: (1) magistrale szeregowo (I²C, SPI, CAN) gdzie aktywność jest paczkowa

– chcemy zarejestrować 100 kolejnych transakcji bez uruchamiania 100 razy ręcznie single-shot, (2) impulsy laserowe, radarowe, RF burst – krótkie zdarzenia o znanej powtarzalności, (3) diagnostyka sporadycznych anomalii – gdy nie wiemy, kiedy się pojawiają, ale wiemy że są rzadkie.

3.7. Pułapki: decymacja w trybie roll i kilka innych

Tryb roll (przewijany) to alternatywny sposób wyświetlania dla bardzo wolnych podstaw czasu – np. 1 s/dz i wolniej. Zamiast czekać aż zbierze się cały rekord (przy 1 GSa/s i 10 dz na ekranie byłoby to 10 sekund), oscyloskop wyświetla próbki na bieżąco, przewijając ekran z prawej do lewej. Pułapka: w trybie roll oscyloskop najczęściej nie zapisuje wszystkich próbek z ADC, tylko co N-tą (decymacja). Sample rate efektywny w trybie roll bywa rzędu kHz nawet jeżeli karta katalogowa deklaruje GSa/s. Konsekwencja: krótkie glitche między próbkami są niewidoczne. Dla diagnostyki glitchów na wolnej podstawie czasu trzeba albo użyć peak detect (jak w sekcji 3.3), albo wyłączyć tryb roll i pracować w normalnym trybie z głębszą pamięcią.

Druga pułapka: sample rate vs pasmo analogowe. Niektóre tanie oscyloskopy mają znacznie wyższy sample rate niż wynikałoby z ich pasma analogowego – widzimy np. „1 GSa/s” przy oscyloskopie o paśmie 50 MHz. Pamiętajmy: pasmo analogowe to twarda granica, którą wyznacza front-end (tłumik, wzmacniacz, ADC).

Trzy ustawienia, od których należy zacząć przy każdym nowym sygnale

✓ **Krok 1** – autoset i ocena widoczności. Przy nowym sygnale klikamy „Autoset” – oscyloskop sam ustawia czułość pionową, podstawę czasu i poziom trigger. Sprawdzamy, czy obraz jest stabilny. Jeżeli tak – przechodzimy do kroku 2. Jeżeli obraz „pływa” – edge trigger nie znajduje stabilnego punktu odniesienia. Wówczas zmniejszamy podstawę czasu (timebase) o jeden bieg, ustawiamy poziom triggera ręcznie blisko punktu połowy amplitudy, włączamy noise reject jeżeli sygnał jest zaszumiony.

✓ **Krok 2** – wybór trybu akwizycji. Jeśli oglądamy szybki sygnał cyfrowy – zostawiamy tryb Sample. Jeśli sygnał jest powolny i może mieć glitche – włączamy peak detect. Jeśli sygnał jest powolny, powtarzalny i chcemy lepszą rozdzielczość pionową – hi-res. Jeśli sygnał jest powtarzalny i szukamy zmiennej, powtarzalnej składowej (np. małego RF na dużym tle szumu) – average z N = 16 lub 64.

✓ **Krok 3** – holdoff jeżeli obraz wciąż pływa. Po ustawieniu trybu akwizycji, jeżeli obraz nadal nie jest stabilny, zwiększamy holdoff – najpierw o czas dwukrotnie krótszy niż podejrzewamy okres sygnału, potem stopniowo modyfikujemy. Większość problemów z migotaniem obrazu da się rozwiązać dwiema zmianami: edge – noise reject, oraz holdoff dłuższy niż wewnętrzna periodyczność sygnału. Dopiero gdy te trzy ustawienia nie wystarczają – warto sięgnąć po wyzwalenie zaawansowane z sekcji 3.5.

Sample rate wyższy niż $5 \times$ pasmo to nadmiar – nie poprawi jakości pomiaru sygnałów bliskich pasma granicznego.

Trzecia pułapka: zbyt długi rekord vs zbyt krótka pamięć. Pamiętamy z sekcji 2.4 wzór: czas obserwacji = pamięć / sample rate. Pokusa zmniejszenia sample rate do uzyskania dłuższego okna jest oczywista, ale: gubimy rozdzielczość czasową (i potencjalnie wpadamy w aliasing!). Lepsze rozwiązania to: głębsza pamięć (kupić oscyloskop z 50...500 Mpts zamiast 1...10 Mpts), albo tryb segmentowy (sekcja 3.6) jeżeli aktywność jest paczkowa.

3.8. Podsumowanie i co dalej

Akwizycja w DSO układa się w spójny ciąg decyzji. Wybieramy sample rate (regułą jest $2,5...5 \times$ pasmo analogowe oscyloskopu, nie pasma podejrzanego sygnału), aby unikać aliasingu. Wybieramy tryb akwizycji odpowiedni do natury sygnału (sample do większości, peak detect do polowania na glitche na długich oknach, hi-res do powolnych sygnałów wysokorozdzielczych, average do redukcji szumu w sygnałach powtarzalnych, envelope do oceny stabilności). Konfigurujemy trigger – w 90% przypadków wystarcza edge trigger z poprawnym sprzężeniem i holdoff, w pozostałych 10% sięgamy po pulse width, runt, time-out lub B-trigger. Wreszcie dla sygnałów paczkowych włączamy tryb segmentowy – pamięć starcza na obserwację okien czasowych rzędu setek razy dłuższych niż przy klasycznej akwizycji.

Karta katalogowa oscyloskopu w obszarze akwizycji powinna podawać: sample rate w trybie jedno- i wielokanałowym (interleaving!), głębokość pamięci w obu trybach, maksymalną liczbę segmentów i rozdzielczość ich time-stampów, listę triggerów wbudowanych (edge, pulse, runt, time-out, sekwencyjny) oraz dekodery protokołarnych (I²C, SPI, UART, CAN, LIN). To są parametry, które w praktyce decydują o tym, jak szybko zlokalizujemy błąd w prototypie – często ważniejsze niż samo pasmo analogowe.

W kolejnym, czwartym rozdziale repetytorium zajdziemy na poziom sondy i front-endu analogowego. Pokażemy, dlaczego pasmo sondy potrafi być wąskim gardłem całego pomiaru, omówimy konstrukcję sond pasywnych 1:10 i 1:1, sond aktywnych, sond różnicowych i sond prądowych, oraz problemy integralności sygnału w pomiarach o paśmie powyżej kilkuset MHz.

Rozdział 4. Sondy i integralność sygnału

W trzech poprzednich rozdziałach omówiliśmy oscyloskop jako sprzęt – jego historię, architekturę wewnętrzną i tryby akwizycji. Czas zająć się drugą połową toru pomiarowego, która równie często decyduje o powodzeniu pomiaru, a której konstruktorzy często nie poświęcają uwagi: sondą. Sonda nie jest pasywnym kawałkiem przewodu



– jest aktywnym uczestnikiem obwodu, który mierzymy. Sonda 1:10 dokłada do obwodu $10 \text{ M}\Omega$ rezystancji równoległej z $10...15 \text{ pF}$ pojemności wejściowej. Aktywna sonda 1:1 ma typowo tylko 1 pF pojemności, ale wymaga zasilania i kosztuje $10...20$ razy więcej. Sonda różnicowa pozwala mierzyć między dwoma punktami nie odniesionymi do masy, ale ma swoje CMRR i właściwe zakresy. Sonda prądowa mierzy bez przerywania obwodu, ale wymaga zrozumienia różnicy między pętlą (cewką) Rogowskiego a cęgami Halla. Cel tego rozdziału jest praktyczny: pokazać, kiedy która sonda jest właściwa, jak unikać typowych pułapek i dlaczego 10-centymetrowy krokodylek masowy potrafi przekłamać pomiar 100 MHz bardziej niż 20% niższe pasmo samego oscyloskopu.

Materiał opieramy na trzech klasycznych źródłach: Tektronix „ABCs of Probes” (dokument 60W-6053, w obecnej rewizji 60W-6053-15) – najbardziej kompletny i wieloletni primer o sondach na rynku; Keysight „Understanding Oscilloscope Probe Specifications” oraz Keysight „Eight Hints for Better Scope Probing” (AN 5989-7894EN) – kombinacja specyfikacji technicznych z praktycznymi wskazówkami; oraz materiały online R&S Essentials. Uzupełniamy je odniesieniami do dokumentów producentów konkretnych sond (Teledyne LeCroy, R&S RT-ZD).

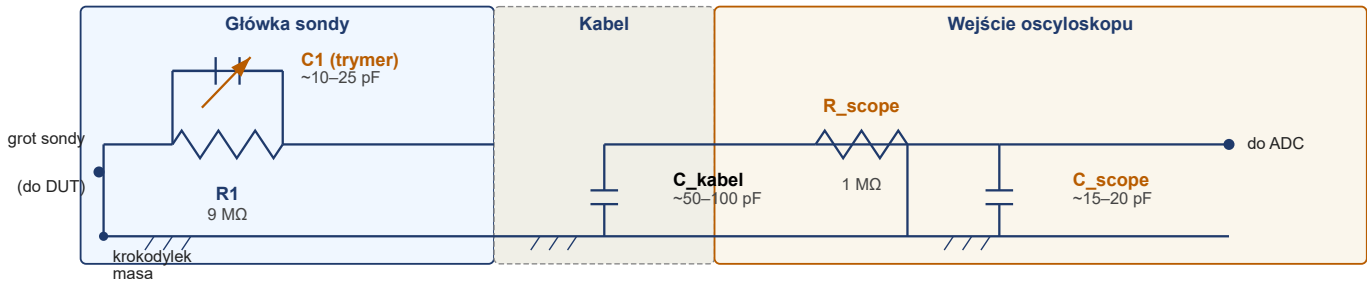
4.1. Sonda jako część toru pomiarowego

Pierwsza zasada probingu: sonda jest częścią obwodu, który mierzymy. Z punktu widzenia elektrycznego obwodu pod testem (DUT – Device Under Test), podłączenie sondy 1:10 to dodanie równoległe do punktu pomiaru jednego rezystora $10 \text{ M}\Omega$ i jednego kondensatora $10...15 \text{ pF}$. Dla wielu sygnałów DC i niskoczęstotliwościowych wpływ jest pomijalny – ale dla szybkich sygnałów i wysokoimpedancyjnych źródeł staje się problemem. Reaktancja 12 pF przy 100 MHz wynosi już ok. 130Ω , co dla obwodu o impedancji źródła $1 \text{ k}\Omega$ oznacza dzielnik tłumiący sygnał o ponad 1 dB . Przy 500 MHz ta reaktancja spada do 27Ω – sonda po prostu zwiera sygnał do masy w paśmie radiowym.

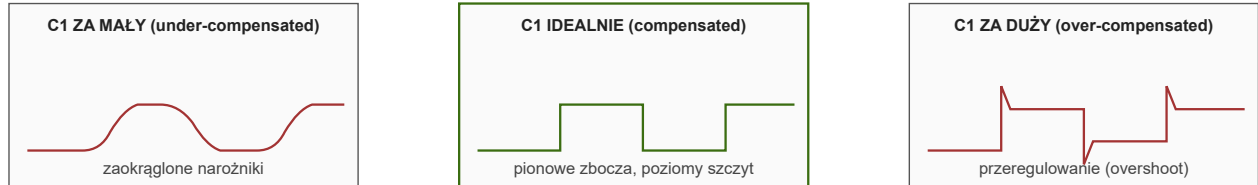
Z tego powodu producenci podają pasmo sondy jako odrębny parametr. Pasma sondy nie jest tożsame z pasmem oscyloskopu – sygnał musi przejść przez oba. Pasma

Schemat zastępczy sondy pasywnej 1:10 z kompensacją

Dlaczego sonda 1:10 ma trymer i co się dzieje, gdy nie jest skompensowana



Co widzi oscyloskop przy kalibrującym sygnale prostokątnym 1 kHz z kalibratora?



Kalibracja sondy (probe compensation):

Sondę 1:10 należy skompensować PRZED każdym pomiarem o paśmie powyżej 1 MHz: podłączyć grot do wyjścia kalibratora oscyloskopu (typowo 1 kHz, 500 mV p-p), pokręcać trymerem C1 śrubokrętem niewielkim ceramicznym aż prostokąt na ekranie ma pionowe zbocza i poziomy szczyt.

Źródło: opracowanie graficzne EP, wzorowane na: Tektronix, „ABCs of Probes” Primer (dokument 60W-6053-15), rozdz. „Passive Probes” oraz „Probe Compensation”; Keysight Technologies, „Becoming Familiar with your Standard Oscilloscope Probe”, Application Note (rozdz. „Probe Compensation”).

Rysunek 12. Schemat zastępczy sondy pasywnej 1:10. Górna część: rezystor R1 (9 MΩ) z równoległym trymerem C1 (~10...25 pF) w głowicy sondy; pojemność kabla C_{kabel} (50...100 pF) i wejście oscyloskopu (R = 1 MΩ, C = 15...20 pF). Dolna część: efekt nieprawidłowej kompensacji widziany na sygnale prostokątnym 1 kHz z kalibratora oscyloskopu. Źródło: opracowanie graficzne EP, wzorowane na: Tektronix, „ABCs of Probes” Primer (dokument 60W-6053-15), rozdz. „Passive Probes” oraz „Probe Compensation”; Keysight Technologies, „Becoming Familiar with your Standard Oscilloscope Probe”, Application Note (rozdz. „Probe Compensation”).

wypadkowe pary sonda – oscyloskop wynika z prostej formuły dwóch filtrów dolnoprzepustowych połączonych szeregowo:

$$\frac{1}{B_{total}^2} = \frac{1}{B_{probe}^2} + \frac{1}{B_{scope}^2}$$

Dla oscyloskopu 500 MHz z sondą pasywną 350 MHz wypadkowe pasmo wynosi tylko ok. 287 MHz – 43% mniej niż deklaracja oscyloskopu. To często niedoceniany szczegół: kupowanie oscyloskopu o paśmie 1 GHz nie ma sensu, jeśli używamy go ze standardowymi sondami 500 MHz dołączonymi w komplecie.

Druga zasada probingu: pętla masy jest częścią sondy. Sonda 1:10 ma standardowo kabelek z krokodylkiem masowym o długości 10...20 cm. Ten kabelek razem z grotem sondy tworzy pętlę, której indukcyjność jest typowo 100...200 nH. W parze z pojemnością wejściową oscyloskopu (~12 pF) tworzy obwód rezonansowy LC o częstotliwości rezonansowej rzędu 100...150 MHz. Każdy szybki proces przejściowy sygnału (zbocze cyfrowe, switching transistor) wzbudza ten obwód i widzimy na ekranie tłumione drgania – ringing (dzwonienie) – który nie jest właściwością mierzonego obwodu, tylko sondy. Pokazujemy to wyraźnie na **rysunku 13**. Mechanizm obronny: spring tip (sprężynka masowa o długości 5 mm zamiast krokodylka 15 cm) redukuje indukcyjność pętli ok. 10-krotnie i przenosi rezonans powyżej 1 GHz, gdzie już nas nie dotyczy.

4.2. Sondy pasywne 1:10 i 1:1 – kompensacja, pojemność, kabel

Sonda pasywna 1:10 to klasyczny dzielnik napięcia: rezystor 9 MΩ w głowicy sondy tworzy razem z wejściem oscyloskopu (1 MΩ) dzielnik 10:1. Ten sam dzielnik widziany dla impedancji wejściowej daje 10 MΩ na DC

Sprawdź kompensację sondy zanim cokolwiek zmierzysz

Każdy oscyloskop ma na panelu czołowym wyjście kalibratora – typowo sygnał prostokątny 1 kHz, 500 mV peak-to-peak. To wyjście istnieje wyłącznie po to, żeby skompensować sondę. Procedura zajmuje 30 sekund i powinna być wykonywana zawsze przy zamianie sondy między oscyloskopami, po wymianie końcówki sondy lub po użyciu sondy w warunkach mechanicznie obciążających.

Krok 1 – podłącz grot sondy do wyjścia kalibratora oscyloskopu, a krokodylek masowy do gniazda masy obok kalibratora. Krok 2 – ustaw na ekranie 3...5 okresów sygnału (ok. 500 μs/dz, czułość 100 mV/dz). Krok 3 – obserwuj kształt zboczy. Krok 4 – jeśli widzisz zaokrąglone narożniki (under-compensated) lub przeregulowanie (over-compensated), znajdź trymer sondy – w nowszych sondach jest to mała śrubka w module kompensacyjnym przy złączu BNC, w starszych w głowicy. Krok 5 – użyj plastikowego lub ceramicznego śrubokrętu (nie metalowego!) i obróć trymer powoli, aż zbocza będą pionowe i szczyt poziomy. Najczęstszy błąd: kompensowanie metalowym śrubokrętem. Metal zwiększa pojemność kompensacji podczas regulacji, więc skompensowana wartość jest niewłaściwa po wyjęciu śrubokręta. Drugi błąd: kompensacja sondy w komplecie z innym oscyloskopem niż ten, do którego będzie podłączona – sondy są kompensowane do konkretnego oscyloskopu, bo C wejściowe różni się między modelami.



– czyli impedancję $10\times$ wyższą niż samego oscyloskopu. Stąd wybór pasywnej 1:10 jako pierwszej sondy: $10\text{ M}\Omega$ to wystarczająco dużo dla większości obwodów cyfrowych i analogowych niskoczęstotliwościowych.

Problem zaczyna się przy częstotliwościach radiowych. Gdyby to był czysto rezystorowy dzielnik, pasmo sondy byłoby ograniczone do kilkudziesięciu kHz przez pojemność kabla – kabel sondy ma typowo $50\ldots 100\text{ pF}$ pojemności do masy, co z rezystorem $9\text{ M}\Omega$ tworzy filtr dolnoprzepustowy o stałej czasowej rzędu kilku mikrosekund. Dlatego w głowicy sondy znajduje się drugi element: kondensator kompensacyjny C_1 równolegle do rezystora $9\text{ M}\Omega$. Wartość C_1 (typowo $10\ldots 25\text{ pF}$) jest dobrana tak, żeby stała czasowa RC w obu gałęziach dzielnika była identyczna. Wtedy dzielnik 10:1 działa dokładnie tak samo dla DC, dla 100 MHz i dla zbroczy nanosekundowych. Schemat zastępczy pokazuje rysunek 12.

Wartość C_1 jest regulowana trymerem (mała śruba w obudowie głowicy lub w kompensowanym module na drugim końcu kabla, blisko BNC). Sondę 1:10 trzeba kompensować przed każdym pomiarem o paśmie powyżej kilkuset kHz – procedura jest opisana w ramce poniżej. Skompensowana źle (over- lub under-compensated) sonda fałszuje kształt zbroczy: under-compensated daje zaokrąglone narożniki sygnału prostokątnego, over-compensated daje przeregulowanie (overshoot) i piki na zbroczach.

Sondy 1:1 (bez tłumienia) są mniej powszechne, ale przydają się dla małych sygnałów. Mają DC impedancję $1\text{ M}\Omega$ (taką samą jak wejście oscyloskopu – sonda 1:1 to praktycznie kabel BNC z grotem). Wady: wyższa pojemność wejściowa ($47\ldots 70\text{ pF}$) i niższe pasmo (typowo $30\ldots 50\text{ MHz}$). Zastosowanie: pomiar bardzo małych sygnałów (mV) przy niskich częstotliwościach, gdy 10-krotne tłumienie sondy pasywnej 1:10 obniża sygnał poniżej szumu oscyloskopu. Niektóre sondy mają przełącznik 1:1/1:10 – ale wartość 1:1 ma wtedy znacznie gorsze pasmo niż 1:10.

4.3. Pasma sondy vs pasmo oscyloskopu

Pasma sondy i pasmo oscyloskopu łączą się jak dwa filtry dolnoprzepustowe szeregowo. Dla pierwszego rzędu

(Gaussian/single-pole roll-off, typowy dla oscyloskopów) wzór jest:

$$f_{total} = \frac{1}{\sqrt{\frac{1}{f_{scope}^2} + \frac{1}{f_{probe}^2}}}$$

Praktycznie: jeśli oscyloskop ma pasmo 500 MHz a sonda 350 MHz , pasmo łączone wynosi 287 MHz – 43% niżej niż deklaracja oscyloskopu. Jeśli sonda ma 1 GHz a oscyloskop 500 MHz , pasmo łączone to 447 MHz – oscyloskop jest wąskim gardłem. Generalna reguła: pasmo sondy powinno być co najmniej $2\times$ wyższe od pasma oscyloskopu, aby sonda nie była wąskim gardłem.

W praktyce na biurku konstruktora typowy zestaw to oscyloskop 200 MHz lub 500 MHz z sondami pasywnymi dołączonymi w komplecie (zwykle o paśmie nominalnie odpowiadającym oscyloskopowi). Producent zazwyczaj specyfikuje pasmo systemu (oscyloskop + sonda) – warto czytać kartę katalogową pod tym kątem. Niektóre oscyloskopy z opcjonalnymi sondami aktywnymi dostają wyższe pasmo (np. Keysight S-series z aktywną InfiniiMax podnosi pasmo z 500 MHz do 4 GHz).

Drugi parametr, na który trzeba patrzeć, to czas narastania sondy (rise time). Wynika on wprost z pasma:

$$t_r[ns] \approx \frac{0,35}{f[GHz]}$$

Dla pasywnej 200 MHz $t_r \approx 1,75\text{ ns}$; dla 500 MHz $\approx 700\text{ ps}$; dla aktywnej 1 GHz $\approx 350\text{ ps}$. Aby mierzyć poprawnie czas narastania sygnału, sonda + oscyloskop muszą mieć co najmniej $5\times$ szybszy czas narastania niż sygnał – inaczej zmierzemy faktycznie czas narastania samego oscyloskopu, nie sygnału.

4.4. Sondy aktywne i różnicowe – kiedy warto

Sondy aktywne mają w głowicy wzmacniacz buforowy (typowo JFET lub bipolarny z bardzo wysoką impedancją wejściową). Najistotniejszą zaletą jest istotna redukcja pojemności wejściowej: z $10\ldots 15\text{ pF}$ typowych pasywnych do 1 pF a nawet $0,6\text{ pF}$ (Keysight InfiniiMax). Konsekwencja: dla pomiaru sygnału cyfrowego o paśmie powyżej kilkuset MHz aktywna sonda nie zniekształca zbroczy ani nie obciąża źródła. Drugą zaletą jest wyższe pasmo – typowo $1\ldots 6\text{ GHz}$, podczas gdy pasywne kończą się na $500\ldots 750\text{ MHz}$.

Koszty aktywnych: cena ($3\text{ 000}\ldots 15\text{ 000}\text{ zł}$ za sondę), konieczność zasilania z oscyloskopu (dedykowany interfejs – TekVPI u Tektroniksa, AutoProbe u Keysighta, ProbeMeter u R&S, ProBus u LeCroy), oraz ograniczone maksymalne napięcie (typowo $\pm 5\text{ V}$ do $\pm 10\text{ V}$ – łatwo uszkodzić aktywną głowicę napięciem 12 V z magistrali zasilania).

Sondy aktywne różnicowe to specjalny przypadek – mają dwa wejścia (plus i minus) zamiast jednego wejścia i krokodylka masowego. Mierzą różnicę napięć między

dwoma punktami, nie muszą być odniesione do masy oscyloskopu. To ma trzy ważne konsekwencje praktyczne. Po pierwsze: można mierzyć między dwoma punktami floating, np. między bramką a źródłem tranzystora MOSFET wysokiej strony w SMPS, gdzie potencjał źródła zmienia się w zakresie setek voltów. Po drugie: wysokie CMRR (typowo 60...80 dB) tłumi sygnały common-mode – jeżeli między oboma punktami pomiaru a masą oscyloskopu jest wspólny szum 1 V, ale sygnał różnicowy to 50 mV, sonda różnicowa wytłumi szum o 1000× i pokaże czysty sygnał. Po trzecie: w wersji wysokiego napięcia (HV differential, np. R&S RT-ZHD15) zakresy do 1500...6000 V – pomiary w SMPS, falownikach, energetyce.

Wady sond różnicowych: CMRR spada z częstotliwością (deklarowane 80 dB jest typowo @ 50 Hz/1 kHz; przy 1 MHz spada do 40 dB; przy 100 MHz może być tylko 20...25 dB). Drugi szczegół: common-mode dynamic range – sonda różnicowa ma ograniczenie nie tylko na wartość różnicową ale i na potencjał wspólny względem masy (common mode). Jeśli próbujemy mierzyć 50 mV różnicy między dwoma punktami przy 800 V common-mode, sonda dostosowana tylko do ±10 V common-mode

po prostu się nasyci. Trzeba dobrać sondę zarówno do zakresu różnicowego, jak i common-mode.

4.5. Sondy prądowe – pętla Rogowskiego, cęgi Halla, sondy AC/DC

Sondy prądowe nie mierzą napięcia, tylko prąd – co wymaga zupełnie innego mechanizmu. Trzy główne technologie: pętla (cewka) Rogowskiego, cęgi z efektem Halla i cęgi z transformatorem prądowym AC. Każda ma swoje zakresy zastosowań i ograniczenia.

Pętla Rogowskiego to giętka, izolowany przewód-elektromagnes, który owijamy wokół testowanego przewodu z prądem. Indukowane napięcie w pętli jest proporcjonalne do pochodnej prądu (dI/dt), więc sonda wymaga integratora (analogowego lub cyfrowego). Zalety: bardzo cienka, łatwo się daje wsunąć w gęsto upakowane miejsca, brak nasycenia magnetycznego nawet przy kiloamperach. Wady: mierzy tylko AC (nie DC), dolne pasmo typowo od kilku Hz, wymaga wyważenia czasu integracji.

Cęgi z efektem Halla mierzą indukcję magnetyczną wokół przewodu prądu – działają na DC i na AC, typowo do 50...100 MHz. To są klasyczne cęgi prądowe stosowane w pomiarach SMPS, falowników, silników BLDC.

Wpływ długości pętli masy na pomiar szybkich stanów nieustalonych

Ten sam sygnał prostokątny 10 MHz mierzony z różnymi przewodami masowymi

A. Sprężynka masowa „spring tip” (~5 mm) – pomiar wierny

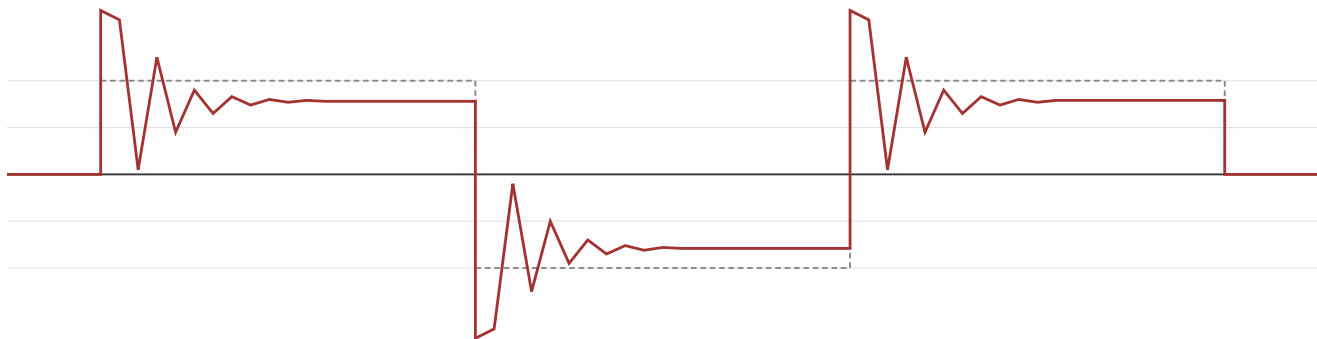
Indukcyjność pętli masy: ~10 nH → częstotliwość rezonansowa zoptymalizowana, brak widocznego ringing.



Zielony przebieg – pomiar; szary przerywany – sygnał idealny (referencja).

B. Pętla masy 15 cm (klasyczny krokodylek na końcu kabla sondy) – dramatyczny ringing

Indukcyjność pętli masy: ~150 nH → rezonans z pojemnością wejściową ~12 pF: $f_{res} = 1/(2\pi\sqrt{LC}) \approx 120$ MHz, długo tłumiona.



Czerwony przebieg – pomiar; szary przerywany – sygnał idealny (referencja).

Każde zbocze sygnału wzbudza tłumione drgania w pętli masy. Czas tłumienia > 100 ns mierzony na ekranie – niemożność oceny realnych zboczy układu.

Wniosek: dla sygnałów o paśmie powyżej 50 MHz długość pętli masy ma większy wpływ na jakość pomiaru niż jakość samej sondy. Używaj spring tip lub ground spring zamiast krokodylka.

Rysunek 13. Wpływ długości pętli masy na pomiar szybkich stanów nieustalonych. Panel A: sonda ze sprężynką masową „spring tip” (~5 mm) – indukcyjność pętli ~10 nH, brak widocznego dzwonienia, pomiar wierny. Panel B: ta sama sonda z krokodylkiem masowym 15 cm – indukcyjność pętli ~150 nH, rezonans LC z pojemnością wejściową oscyloskopu (~12 pF) przy ok. 120 MHz, widoczne tłumione drgania na każdym zboczu sygnału. Źródło: opracowanie graficzne EP, wzorowane na: Tektronix, „ABCs of Probes” Primer (60W-6053-15), rozdz. „Grounding Best Practices”; Keysight, „Eight Hints for Better Scope Probing” (AN 5989-7894EN)



Wymagają zasilania (z baterii lub z interfejsu sondy). Dokładność typowo 1...3% odczytu, offset DC bywa problemem (trzeba używać funkcji zerowania przed pomiarem). Zakresy od ± 30 A do ± 1000 A w zależności od modelu.

Cęgi z transformatorem prądowym AC są najtańsze, ale działają tylko na AC – typowo od kilkudziesięciu Hz do kilkuset kHz. Stosowane w pomiarach mocy sieci, prądu wejściowego zasilaczy, prądów w transformatorach. Najczęściej spotyka się je jako akcesoria

do multimetrów i analizatorów mocy, rzadziej do oscyloskopów. W ofercie oscyloskopowej dominują obecnie cęgi Halla AC/DC.

Pułapka praktyczna: sondy prądowe mają specyfikację „skin effect derating” – przy wysokich częstotliwościach prąd płynie tylko w warstwie naskórkowej przewodnika, sonda mierzy go z błędem zależnym od grubości przewodu. Druga pułapka: zerowanie offsetu DC przed pomiarem musi być wykonane bez prądu w przewodzie (pętla otwarta), nie obok przewodu z prądem.

4.6. Pętla masy – najczęstszy błąd w pomiarach RF

Po kompensacji sondy najczęstszym powodem złych pomiarów jest długi przewód masowy. Krokodylek na końcu kabla 15 cm tworzy pętlę o indukcyjności typowo 100–200 nH. W parze z pojemnością wejściową oscyloskopu (12 pF) i pojemnością sondy (10 pF łącznie z kablem) tworzy obwód rezonansowy LC o częstotliwości rezonansowej rzędu 100...150 MHz.

Każde szybkie zbocze (typowo czas narastania < 10 ns) ma w swoim widmie składowe sięgające powyżej 50 MHz,

Cztery typy sond oscyloskopowych – kluczowe parametry i zastosowania

Pasywna 1:10, aktywna jednokońcowa, aktywna różnicowa, prądowa

	Pasywna 1:10	Aktywna jednokońcowa	Aktywna różnicowa	Prądowa
Pasmo typowe	200 MHz – 500 MHz	1 GHz – 6 GHz	100 MHz – 2 GHz	DC – 100 MHz (cęgi)
Impedancja wejściowa	10 M Ω 10–15 pF spadek z f (capacitive load)	100 k Ω – 1 M Ω ~1 pF znacznie niższe obciążenie	1 M Ω ~1–4 pF niska, symetryczna	indukcyjne sprzężenie (niebezpośrednie)
Maks. napięcie	300–600 V (CAT II/III)	± 5 V do ± 10 V delikatna, łatwo uszkodzić	do 1500 V (HV diff) odsprężnięta od masy	do 1000 A (zależy od cęgi) pętla Rogowski / Hall
CMRR @ 1 MHz	n.d. (jednokońcowa)	n.d. (jednokońcowa)	60–80 dB izolacja masy	izolacja galwaniczna
Wymaga zasilania	nie	tak (z oscyloskopu)	tak (z oscyloskopu)	cęgi AC: nie; Hall: tak
Cena (orient.)	200–800 zł	3 000 – 15 000 zł	5 000 – 30 000 zł	2 000 – 20 000 zł

Typowe zastosowania

- cyfrowe sygnały 3,3 V / 5 V - zegary mikrokontrolerów - magistrale szeregowo niskie (UART, I²C, SPI) 90% codziennych pomiarów na biurku konstruktora	- szybkie magistrale danych (USB 2.0/3.0, HDMI) - zegary > 200 MHz - sygnały z impedancją źródła > 1 kΩ - gdy 10 pF wejścia pasywnej zakłóca obwód mierzony	- pomiary między dwoma punktami floating - SMPS – sygnały bramki/source tranzystora wysokiej strony - napięcia różnicowe USB/CAN/LIN/Ethernet - pomiary RF nie referowane do masy	- prąd przez tranzystor SMPS - prąd wejściowy zasilacza - prąd w cewce silnika BLDC - prąd uzwojenia transformatora - bez konieczności przerywania obwodu (cęgi otwarte) - korzystanie z funkcji pomiaru mocy w oscyloskopach
--	--	--	--

Zasada doboru:

zaczynamy od pasywnej 1:10 (taniej, niezniszczalnej, wystarczającej w 90% przypadków); aktywne, różnicowe lub prądowe – tylko gdy pasywna nie wystarcza.

Źródło: opracowanie graficzne EP, wzorowane na: Tektronix „ABCs of Probes” Primer (60W-6053-15), tab. „Probe Comparison”; Keysight, „Eight Hints for Better Scope Probing” (AN 5989-7894EN); R&S „Types of oscilloscope probes” (online tutorial).

Rysunek 14. Porównanie czterech głównych typów sond oscyloskopowych. Pasywna 1:10 to standard „z pudełka” do większości pomiarów. Aktywna jednokońcowa: dla wysokich częstotliwości i wysokoimpedancyjnych źródeł. Aktywna różnicowa: dla pomiarów floating i z wysokim CMRR. Prądowa: dla pomiarów prądu bez przerywania obwodu. Źródło: opracowanie graficzne EP, wzorowane na: Tektronix „ABCs of Probes” Primer (60W-6053-15), tab. „Probe Comparison”; Keysight, „Eight Hints for Better Scope Probing” (AN 5989-7894EN); R&S „Types of oscilloscope probes” (online tutorial)

które wzbudzają ten rezonans. Widzimy na ekranie tłumione drgania (ringing) – które nie są właściwością mierzonego sygnału, tylko sondy. Pokazujemy to wyraźnie na **rysunku 13**: ten sam sygnał prostokątny mierzony z krótką sprężynką masową (5 mm, ok. 10 nH) i z klasycznym krokodylkiem (15 cm, ok. 150 nH) daje dwa zupełnie różne obrazy.

Mechanizm obronny: spring tip albo ground spring. To metalowa sprężynka w kształcie „U” lub spirali, długości 5...10 mm, która zastępuje krokodylkowy kabel masowy. Mocujemy ją bezpośrednio do koszulki BNC sondy. Indukcyjność zredukowana ok. 10-krotnie, częstotliwość rezonansowa obwodu LC przesuwana powyżej 1 GHz, gdzie szybkie zbocza już nie wzbudzają rezonansu. Większość sond aktywnych ma spring tip w komplecie; do sond pasywnych trzeba ją dokupić osobno (typowo 30–100 zł za zestaw kilku sztuk).

Drugi częsty błąd to wiele różnych pętli masy w pomiarach kilku-kanalowych. Każdy kanał oscyloskopu ma swój krokodylek masowy, więc czterokanałowy pomiar oznacza cztery różne pętle masy z DUT do oscyloskopu. Te pętle mogą tworzyć ground loops, których szum jest dodawany do każdego pomiaru. Mechanizm obronny: w pomiarach wysokoczęstotliwościowych podłączamy wszystkie masy do tego samego punktu DUT (single-point grounding).

4.7. Cztery typy sond w jednej tabeli – którą wybrać

Na podstawie rozważań sekcji 4.1...4.5 możemy podsumować cztery główne typy sond i ich zastosowania. Tabela na **rysunku 14** przedstawia kluczowe parametry – pasmo, impedancję wejściową, maksymalne napięcie, CMRR, wymagania zasilania i orientacyjne ceny – oraz typowe zastosowania każdej z nich.

Generalna zasada doboru: zaczynamy od pasywnej 1:10 jako pierwszej sondy – jest tania, niezniszczalna, nie wymaga zasilania i wystarcza w 90% przypadków pomiarów na biurku konstruktora. Aktywna, różnicowa lub prądowa wchodzi do gry dopiero wtedy, gdy pasywna nie wystarcza: aktywna gdy mierzymy >500 MHz lub źródło o wysokiej impedancji, różnicowa gdy punkty pomiaru są floating względem masy lub gdy mamy duży szum common-mode, prądowa gdy musimy mierzyć prąd bez przerywania obwodu.

4.8. Pułapki: napięcia wysokie, sondy izolowane, bezpieczeństwo CAT

Każda sonda ma w karcie katalogowej dwa parametry napięciowe: maksymalne napięcie wejściowe (typowo deklarowane dla DC + AC peak) oraz kategorię pomiarową (CAT II/CAT III/CAT IV). Kategoria pomiarowa jest

Kiedy potrzebujesz sondy izolowanej, a kiedy nie

Sonda izolowana (HV differential lub IsoVu) jest konieczna w trzech sytuacjach. Pierwsza: pomiar między dwoma punktami obwodu, z których ŻADEN nie jest masą oscyloskopu – np. bramka tranzystora MOSFET wysokiej strony w mostku H, gdzie źródło pływa w zakresie 0...400 V. Druga: pomiar sygnału AC na poziomie sieci zasilającej (230 V/400 V) lub wyższym. Trzecia: pomiar w obwodzie z izolacją galwaniczną, której nie chcemy „zszyc” podłączając krokodylek masowy. Sonda zwykła pasywna 1:10 wystarczy w obwodach niskoenergetycznych, gdzie pomiar jest robiony względem masy obwodu (np. sygnał ground płytki) i obwód ma izolację od sieci zasilającej – typowo wszystkie obwody zasilane z USB, z baterii lub z izolowanego zasilacza laboratoryjnego. Najgorszy wybór: brak sondy izolowanej tam, gdzie jest potrzebna, i próba ‘obejścia’ przez transformator izolujący 1:1 podpięty do oscyloskopu. To rozwiązanie usuwa uziemienie PE oscyloskopu – w razie błędu pomiaru lub awarii izolacja oscyloskopu spada na operatora. Wszyscy producenci oscyloskopów wyraźnie odradzają tę technikę w manualach bezpieczeństwa.

standardem IEC 61010 i określa, w jakich obwodach sonda może być bezpiecznie używana. CAT II to instalacje domowe poza tablicą rozdzielczą (typowe pomiary do 600 V), CAT III to instalacje przed tablicą rozdzielczą (typowo 600 V), CAT IV to bezpośrednie wejście z sieci dystrybucyjnej (typowo 600 V z marginesem na przebiegi przejściowe).

Klasyczna pułapka: zwykła sonda pasywna 1:10 nie ma izolacji od masy oscyloskopu – jej krokodylek masowy łączy się bezpośrednio z masą oscyloskopu, która łączy się z PE w przewodzie zasilającym oscyloskop. Próba pomiaru bezpośrednio na zaciskach 230 V przy podłączonym do tej samej sieci oscyloskopie kończy się zwarcie fazy z PE przez sondę – spalone bezpieczniki w najlepszym przypadku, pożar w gorszym. Do pomiarów AC sieciowych obowiązkowa jest sonda różnicowa wysokiego napięcia (HV differential) z odpowiednim CAT rating.

Druga pułapka: oscyloskop podłączony do sieci zasilającej z transformatorem izolującym (1:1, często stosowany przez serwisantów) nie jest bezpieczny do pomiarów pod napięciem. Transformator izoluje masę oscyloskopu od PE sieci, ale wciąż ma pojemność izolacji rzędu pF, która przy wysokim dV/dt może spowodować przepływ prądu zakłócającego pomiar. Sonda izolowana (np. R&S RT-ZISO, Tektronix IsoVu) ma izolację optyczną wewnątrz sondy, więc nawet kilowoltowe napięcie wspólne (common-mode) nie wpływa na pomiar. Sondy IsoVu są drogie (10 000...30 000 zł), ale dla diagnostyki SMPS wysokiego napięcia są niezastąpione.

Trzecia pułapka: sondy z deklaracją wysokiego napięcia (np. R&S RT-ZP10 do 600 V CAT II) mają ten rating tylko dla sondy jako całości – ale krokodylek masowy w typowej długości 15 cm jest słabym ogniwnem izolacyjnym. Producenci sond wysokiego napięcia dołączają w komplecie alternatywne akcesoria masowe – dwa kabelki o różnej długości (typowo 6 cali/15 cm oraz 18 cali/46 cm) oraz sprężynę masową (ground spring) o izolacji odpowiadającej



kategori napięciowej sondy. Czytanie karty katalogowej sondy do końca, włącznie z „Safety Information”, jest obowiązkowe przed pomiarem na napięciu.

4.9. Podsumowanie i co dalej

Sonda jest pełnoprawnym elementem toru pomiarowego, nie pasywnym kawałkiem przewodu. Dobór sondy

do zadania – pasywna 1:10 do większości pomiarów, aktywna jedno- lub różnicowa do wysokich częstotliwości i pomiarów floating, prądowa do pomiaru prądu bez przerywania obwodu – powinien być świadomy i poprzedzony analizą pasma, impedancji źródła, zakresu napięcia i kategorii bezpieczeństwa. Pasma łączne sondy i oscyloskopu daje filtr szeregowy o paśmie wypadkowym niższym od mniejszego z dwóch – warto pamiętać tę formułę przy kupowaniu nowego oscyloskopu. Kompensacja sondy 1:10 trwa 30 sekund i jest obowiązkowa – źle skompensowana sonda fałszuje zbocza wszystkich pomiarów wysokoczęstotliwościowych. Pętla masy jest typowo gorszym ograniczeniem pasma niż sama sonda – spring tip zamiast krokodyłka redukuje ringing o rząd wielkości.

W kolejnym, piątym rozdziale repetytorium przyjrzymy się trendowi, który najmocniej zmienił ofertę oscyloskopów w ostatnich 5 latach: wejściu 12-bitowych ADC i wbudowanych kanałów logicznych (MSO) do klasy budżetowej, czyli poniżej 10 000 zł. Pokażemy, jak Siglent, R&S, UNI-T i Rigol zmienili rynek między 2020 a 2026, co to oznacza dla konstruktora i które kompromisy konstrukcyjne zostały wymuszone, by 12-bitowy ADC trafił na biurko każdego warsztatowca.

WM

Bibliografia

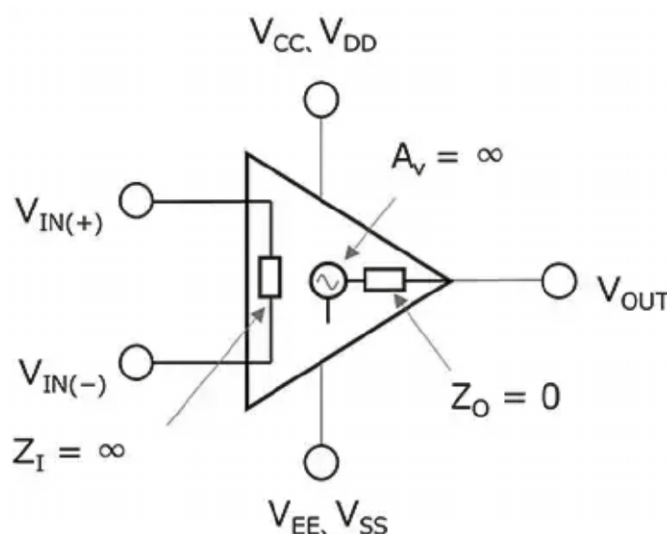
- [1] Keysight Technologies, „Basic Oscilloscope Fundamentals”, Application Note 5989-8064EN (publikacja 2014, ostatnia rewizja: marzec 2025).
- [2] Tektronix, „XYZs of Oscilloscopes – Primer”, dokument 03W-8605 (kolejne wydania od 2000 r.; wydanie 03W-8605-4: 2009 r.; aktualne wydanie online dostęp: kwiecień 2026).
- [3] Rohde & Schwarz, „Oscilloscope Fundamentals/Oscilloscope Basics” (materiał edukacyjny producenta, dostęp: kwiecień 2026).
- [4] Keysight Technologies, „Evaluating Oscilloscope Sample Rate vs. Sampling Fidelity”, Application Note 5989-5732 (pierwotne wydanie 2011, ostatnia rewizja 2023).
- [5] Keysight Technologies, „Segmented Memory Acquisition for InfiniiVision Series Oscilloscopes”, Data Sheet 5989-7833EN (rewizja grudzień 2023). URL: keysight.com/us/en/assets/7018-01736/data-sheets/5989-7833.pdf
- [6] Keysight Technologies, „Using a Scope’s Segmented Memory to Capture Signals More Efficiently”, Application Note 5989-4932EN (rewizja marzec 2024). URL: keysight.com/us/en/assets/7018-01382/application-notes/5989-4932.pdf
- [7] Tektronix, „Using FastFrame™ Segmented Memory”, Application Note 55W-12112 (klasyczne opracowanie dla rodziny DPO7000; wydanie 55W-12112-2). URL: download.tek.com/document/55W_12112_2.pdf
- [8] Tektronix, „Using FastFrame™ Segmented Memory on the 4/5/6 Series MSO”, Technical Brief 55W-61299 (najnowsza dostępna rewizja: 55W-61299-2; dla obecnych serii MSO 4/5/6). URL: download.tek.com/document/FastFrame_Tech-Brief_55W-61299-2.pdf
- [9] Rohde & Schwarz, „Oscilloscope Basics” (primer producenta; rozdział 7 poświęcony w całości wyzwalaniu – minimum detectable glitch, threshold, jitter wyzwalania, edge trigger, glitch, runt, mask violation; dostęp: kwiecień 2026). URL: instrumentationlab.berkeley.edu/sites/default/files/Articles/Oscilloscope-Basics.pdf
- [10] Tektronix, „ABCs of Probes – Primer”, dokument 60W-6053 (klasyczny tutorial sond Tektroniksa; pierwsze wydanie 1990; najnowsza dostępna rewizja 60W-6053-15). URL: download.tek.com/document/ABCs%20of%20Probes%2060W-6053-15.pdf
- [11] Keysight Technologies, „Understanding Oscilloscope Probe Specifications”, Application Note (przegląd parametrów sond: capacitive loading, CMRR, attenuation; dostęp: kwiecień 2026). URL: keysight.com/us/en/assets/3124-1574/application-notes/Understanding-Oscilloscope-Probe-Specifications.pdf
- [12] Keysight Technologies, „Eight Hints for Better Scope Probing”, Application Note 5989-7894EN (praktyczne wskazówki probing dla konstruktorów; dostęp: kwiecień 2026).
- [13] Rohde & Schwarz, „Types of oscilloscope probes” oraz „Oscilloscope probe tips and how to use them”, R&S Essentials (materiał edukacyjny producenta online; cztery typy sond, procedura kompensacji; dostęp: kwiecień 2026). URL: rohde-schwarz.com/us/products/test-and-measurement/essentials-test-equipment/rs-essentials-digital-oscilloscopes/types-of-oscilloscope-probes_257575.html
- [14] Keysight Technologies, „Becoming Familiar with your Standard Oscilloscope Probe”, Application Note (procedura kompensacji sondy 1:10; dostęp: kwiecień 2026).

Wzmacniacze operacyjne i komparatory

Każdy, kto uczył się elektroniki analogowej, na pewno zetknął się ze słynnym „magicznym trójkątem” – przyrządem o nieskończenie wielkim wzmocnieniu różnicowym, zerowej impedancji wyjściowej i takie tam... W praktycznych warunkach nie ma możliwości zakupu owego „idealnego” podzespołu, ale można dobrać taki, który będzie najlepszy z możliwych. Czemu komparatory i wzmacniacze operacyjne to prawie to samo, ale nie do końca, dlaczego celowo ograniczamy ich pasmo przenoszenia, czemu kompromisy są potrzebne i co z tym wszystkim ma wspólnego żarówka? Zapraszam do lektury!

Wzmacniacze operacyjne, bo one będą gwiazdami tego artykułu, wzięły swoją nazwę od dosyć nieoczywistego dzisiaj zastosowania – od wykonywania operacji matematycznych w komputerach analogowych. Trzeba przyznać, że ich parametry (w wersji idealnej) powodują, że stają się wręcz stworzone do odejmowania, dodawania, całkowania i różniczkowania sygnałów. Owszem, dzisiaj nadal używamy układów całkujących i różniczkujących, jak również sumatorów i subtraktorów, jednak rzadko myślimy o nich jako o układach realizujących działania (inaczej: operacje) matematyczne. To tyle tytułem wyjaśnienia ich nazwy.

Schemat blokowy idealnego wzmacniacza operacyjnego znajduje się na **rysunku 1**. Ponieważ mówimy o ideałach, to jego impedancja wejściowa powinna być nieskończenie wielka, co oznacza, że przez wejścia nie płyną jakiegokolwiek prądy. Wzmocnienie różnicowe powinno być, co nieszczerólnie dziwi, nieskończenie wysokie. Z kolei impedancja wyjściowa ma wynosić $0\ \Omega$ – tak, okrągłe zero. Ze swojej strony dodam jeszcze, co rzadko się uwzględnia przy wymienianiu cech wzmacniacza idealnego, że jego pole wzmocnienia, definiowane jako GBP (Gain-Bandwidth Product – częstotliwość, przy której wzmocnienie spada do $1\ \text{V/V}$) powinno być nieskończenie szerokie ($\text{GBP} \rightarrow \infty$), a czas narastania nieskończenie krótki ($t_r \rightarrow 0$). Ponadto, współczynnik tłumienia tętnień zasilania (PSRR) oraz współczynnik wzmocnienia sumacyjnego (CMRR) powinny dążyć do zera. No, tośmy się pośmiali... Idealnych



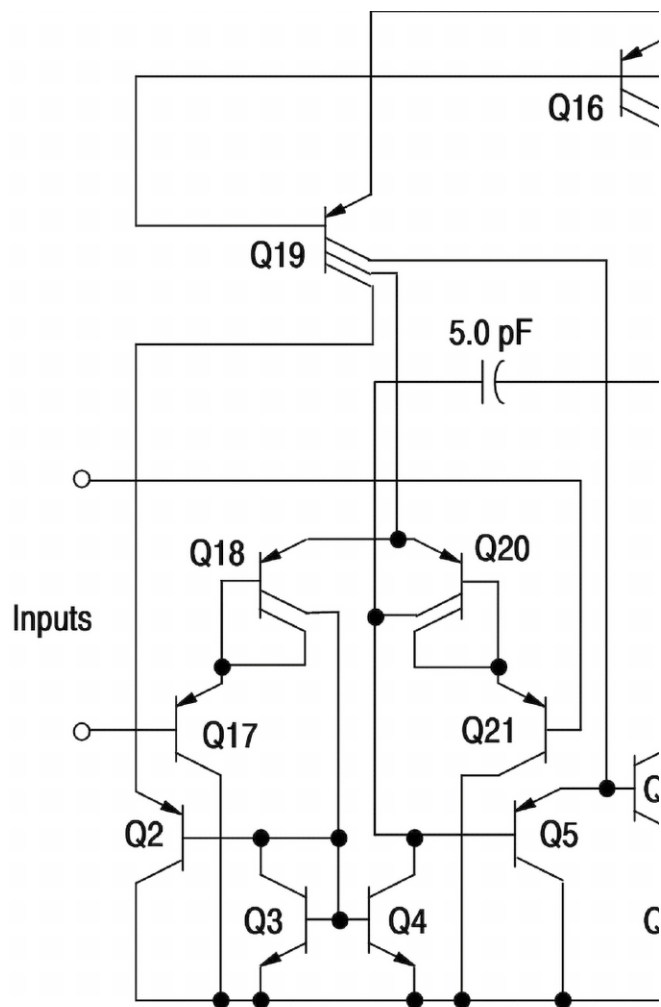
Rysunek 1. Schemat blokowy idealnego wzmacniacza operacyjnego [1]

podzespołów jednak nie ma, a im bardziej złożony jest element, tym więcej cech odróżnia go od tego ideału. Jak pokażę w dalszej części artykułu, są jednak możliwe pewne kompromisy, które pod pewnymi względami dają nam coś, co możemy uznawać za dużą dozę doskonałości.

Bardzo często obok wzmacniaczy operacyjnych występują, choć dosyć nieśmiało, ich bliscy kuzyni, czyli komparatory. Nie ma się co dziwić, bo od strony układowej wyglądają niemal tak samo, zaś na schematach wręcz identycznie, a jednak ich przeznaczenie jest zdecydowanie mniej różnorodne. Podobnie jak wzmacniacze operacyjne, komparatory mają za zadanie wzmocnienie różnicy między dwoma potencjałami przyłożonymi do ich wejść, lecz tylko po to, by zasygnalizować, który z nich ma większą wartość. I to tyle, koniec romantycznej opowieści. Mimo tego, komparatory są bardzo istotne w nowoczesnych układach analogowych, a już na pewno stanowią doskonały „pomost” na styku świata analogowego z cyfrowym.

Stopień wejściowy

Zarówno wzmacniacze operacyjne, jak i komparatory, zawsze posiadają wejście różnicowe. I zawsze jest ono realizowane przy użyciu mniej lub bardziej rozbudowanego układu różnicowego. Przykładowy, znajdujący się w strukturze bardzo popularnego układu LM358, można zobaczyć na **rysunku 2**. Podstawowy układ różnicowy (transystory Q18 i Q20) są sterowane przez wtórnik napięciowy



Rysunek 2. Schemat ideowy stopnia wejściowego układów LM358 i pokrewnych [2]

zrealizowane na tranzystorach Q17 i Q21 – wszystkie o polaryzacji PNP. Daje to dwie informacje, użyteczne w praktyce:

- wejście będzie dobrze obsługiwało potencjały bliskie ujemnej linii zasilającej;
- prądy polaryzujące tranzystory stopnia wejściowego będą wypływały z wyprowadzeń wejściowych, a nie do nich wpływały.

Na temat pierwszej informacji z reguły szerzej rozpisuje się każda karta katalogowa w swojej dalszej części, w których podany jest zakres obsługiwanych napięć, ale rzut oka na schemat daje nam pewną intuicję w tym zakresie: niskie potencjały owszem, wysokie już nie za bardzo. Z kolei druga informacja również ma swoje odzwierciedlenie w tabelach w postaci ujemnych wartości prądów wejściowych (rysunek 3), ale nie zawsze bywa to przestrzegane przez producentów. Ma to dla nas, konstruktorów elektroników, niebagatelne znaczenie w aplikacjach wymagających przenoszenia składowej stałej lub tam, gdzie trzeba na nią uważać. Otóż dodanie rezystancji szeregowej z wejściem, z którego wypływa prąd spowoduje podniesienie potencjału tego wejścia, co ma bezpośrednie przełożenie na napięcie wyjściowe: przy wejściu nieodwracającym spowoduje jego podniesienie, zaś przy odwracającym obniżenie.

Piszę o tym nie bez przyczyny, ponieważ wiele lat temu realizowałem prosty układ wzmacniacza (nieodwracającego) sygnał z czujnika dymu. Nie skompensowałem odpowiednio rezystancji „widzianych” przez wejścia i okazało się, że napięcie wyjściowe jest podniesione o jakiś niewytłumaczalny offset, zawsze z konkretnym znakiem. Jak się okazało, winna był właśnie prąd „wypływający” z wejść. Dobór odpowiednich rezystancji załatwił temat.

Przykład z rysunku 2 jest bardzo często spotykany, by nie powiedzieć – banalny. Na rysunku 4 znajduje się schemat stopnia wejściowego układu, który można postawić obok ideału (pod względem wartości prądów wejściowych) i niewiele się pomylić. Tak, szumi. Tak, ma offset napięciowy jak z Grudziądza do Honolulu. Ale tranzystory JFET z kanałem typu P w stopniu wejściowym dają fenomenalną rezystancję wejściową, rzędu 1 TΩ, przy prądzie wejściowym rzędu 50...400 pA (maksymalnie 8 nA – w podwyższonej temperaturze zaczyna się ujawniać prąd wynikający z samoistnej generacji termicznej par elektron-dziura). W zakresie temperatur „pokojowych” i niższych w takim układzie nie trzeba szczególnie przejmować się kompensacją.

Characteristic	Symbol	LM258			LM358, LM358E			LM358A			Unit
		Min	Typ	Max	Min	Typ	Max	Min	Typ	Max	
Input Offset Voltage $V_{CC} = 5.0 \text{ V to } 30 \text{ V}$, $V_{IC} = 0 \text{ V to } V_{CC} - 1.7 \text{ V}$, $V_O \approx 1.4 \text{ V}$, $R_S = 0 \Omega$ $T_A = 25^\circ\text{C}$ $T_A = T_{\text{high}}$ (Note 4) $T_A = T_{\text{low}}$ (Note 4)	V_{IO}	–	2.0	5.0	–	2.0	7.0	–	2.0	3.0	mV
Average Temperature Coefficient of Input Offset Voltage $T_A = T_{\text{high}}$ to T_{low} (Note 4)	$\Delta V_{IO}/\Delta T$	–	7.0	–	–	7.0	–	–	7.0	–	$\mu\text{V}/^\circ\text{C}$
Input Offset Current $T_A = T_{\text{high}}$ to T_{low} (Note 4)	I_{IO}	–	3.0	30	–	5.0	50	–	5.0	30	nA
Input Bias Current $T_A = T_{\text{high}}$ to T_{low} (Note 4)	I_{IB}	–	–45	–150	–	–45	–250	–	–45	–100	
		–	–50	–300	–	–50	–500	–	–50	–200	

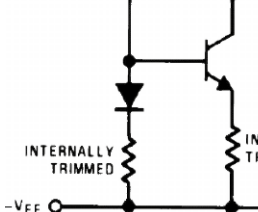
Rysunek 3. Ujemne wartości prądów wejściowych układów LM358 i pokrewnych [2]

cją prądów polaryzujących wejścia, co jest wygodne w aplikacjach np. ze zmiennym wzmacnieniem.

No dobrze – ktoś mógłby powiedzieć – czym się zachwycać, skoro mamy do dyspozycji wzmacniacze operacyjne CMOS, posiadające prąd wejściowy rzędu femtoamperów czy pojedynczych pikoamperów, jak chociażby popularny OPA336 (rysunek 5). Wszak izolator podbramkowy zawsze będzie lepszym... właśnie, izolato-

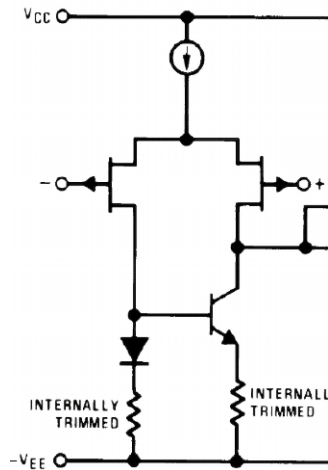
rem, niż spolaryzowane zaporowo złącze PN. To dopiero jest ideał, wejścia prezentują sobą prawdziwe rozwarci!

Owszem, wejścia tak, tylko z reguły cierpi reszta układu: niskie napięcie zasilania, wysokie szумы, bardzo niski slew-rate lub inne „bólączki”. Za to pobór prądu przez wzmacniacz operacyjny też z reguły jest najprzyjemniejszy w wydaniu układów CMOS, bowiem takiemu OPA336 wystarczy 20 μA , w porywach do 42 μA [4]. Nie jest prawdą, że wzmacniacze CMOS muszą być wolne – pierwszy z brzegu OPA356 ma wejścia CMOS a cechuje się GBW sięgającym 200 MHz [5] i slew-rate rzędu $-360...+300 \text{ V}/\mu\text{s}$. (co się dziwić, szerokie pasmo) powodują, że układ ten bynajmniej nie jest wolny. Ale napięcie zasilające nie większe niż 5,5 V powoduje, że układ ten nie może występować w każdym zastosowaniu, choć z pewnością jest kuszący.



Rysunek 4. Schemat ideowy stopnia wejściowego układu TL082 [3]

Nie wolno zapominać o wzmacniaczach rail-to-rail, które mają stopień wejściowy zbudowany inaczej niż w poprzednio omówionych układach. Zamiast pojedynczego układu różnicowego, zawiera on tak naprawdę dwa układy różnicowe, zbudowane z komplementarnych tranzystorów. I tak oto układy różnicowe z tranzystorami NPN (lub z kanałem typu N) dobrze obsługują wysokie napięcia, bliskie



Rysunek 4. Schemat ideowy stopnia wejściowego układu TL082 [3]

dotychczasowej linii zasilającej, z kolei układy różnicowe z tranzystorami PNP (lub z kanałem typu P) sprawdzają się w obsłudze napięć niskich, bliskich ujemnej linii zasilającej. W dalszej części wzmacniacza operacyjnego trzeba połączyć sygnały wychodzące z obu tych wzmacniaczy, co samo w sobie jest bardzo ciekawe. Pewną zaletą jest wzajemna kompensacja prądów polaryzujących bazy tranzystorów bipolarnych przez tranzystory o odmiennym typie polaryzacji. Nie darzę zbytnim zaufaniem wzmacniaczy rail-to-rail, ponieważ w pobliżu owych rails są one z reguły bardzo wolne, by nie rzec: zdegenerowane, a i to przy odpowiednio wysokiej rezystancji obciążenia. Traktuję je jako ostateczność, ponadto wiele dzisiejszych wzmacniaczy CMOS zachowuje się jak wzmacniacze rail-to-rail przy wysokiej rezystancji obciążenia. Wole w swoich układach dążyć do uzyskania pewnego „naddatku” napięcia, który zapewni prawidłową pracę obwodów wzmacniacza operacyjnego – uważam to za pewniejsze niż liczenie, że napięcie wyjściowe naprawdę sięgnie linii zasilającej, co jest szczególnie trudne przy większym poborze prądu z wyjścia wzmacniacza operacyjnego.

Nowoczesne wzmacniacze operacyjne nie kończą się bynajmniej na opisanych tutaj układach. Ich oferta u popularnych dystrybutorów liczy tysiące pozycji, więc opisanie ich wszystkich jest zajęciem karkołomnym. Zwrócę uwagę na ciekawy układ: LT1364 (podwójny) tudzież LT1365 (poczwórny) wzmacniacz operacyjny o slew-rate, uważa... 1000 V/ μ s! Mają szeroki zakres napięć zasilających ($\pm 2,5 \dots 15$ V) i stosunkowo niski pobór prądu, rzędu 6...8 mA na wzmacniacz [7]. Bardzo dobrze sprawują się w roli driverów kabli koncentrycznych 50 Ω i 75 Ω . To wzmacniacz operacyjny o dosyć unikalnej technologii, w której prąd stopnia wejściowego jest zmieniany dynamicznie w zależności od napięcia różnicowego. Przy dużych skokach, do obsługi których których przydaje się wysoki slew-rate, prąd stopnia wejściowego jest zwiększany, zaś przy małych sygnałach pracuje on z niewielkim prądem. To skutkuje dwiema przeciwstawnymi cechami: niskim poziomem szumów (w obrębie słabych sygnałów)

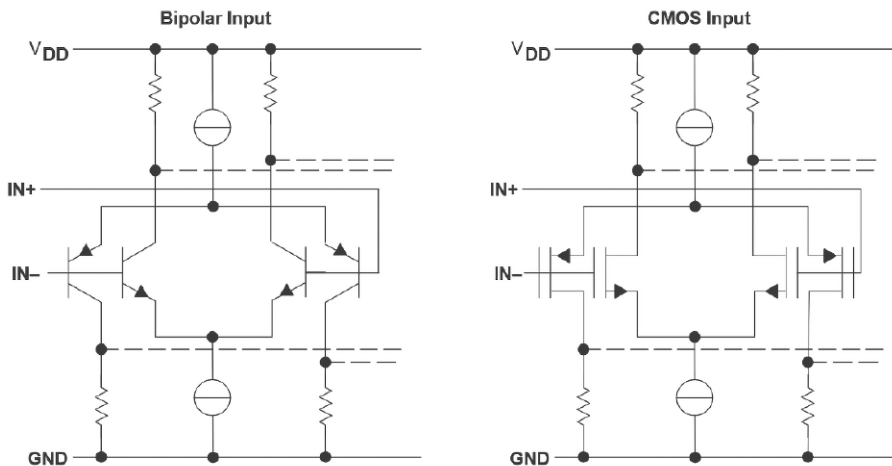
PARAMETER	CONDITION	OPA336N, U OPA2336E, P, U			OPA336NA, UA OPA2336EA, PA, UA OPA4336EA			OPA336NJ, UJ			UNITS
		MIN	TYP ⁽¹⁾	MAX	MIN	TYP	MAX	MIN	TYP	MAX	
OFFSET VOLTAGE Input Offset Voltage vs Temperature vs Power Supply Over Temperature Channel Separation, dc	V_{OS} dV_{OS}/dT PSRR $V_S = 2.3V$ to $5.5V$ $V_S = 2.3V$ to $5.5V$		± 60 ± 1.5 25 0.1	± 125 100 130		* * * *	± 500 * * *		± 500 * * *	± 2500 * * *	μV $\mu V/^{\circ}C$ $\mu V/V$ $\mu V/V$ $\mu V/V$
INPUT BIAS CURRENT Input Bias Current Over Temperature Input Offset Current	I_B I_{OS}		± 1 ± 1	± 10 ± 60 ± 10		* * *	* * *		* * *	* * *	pA pA pA

Rysunek 5. Prądy wejściowe układów OPA336 i pokrewnych [4]

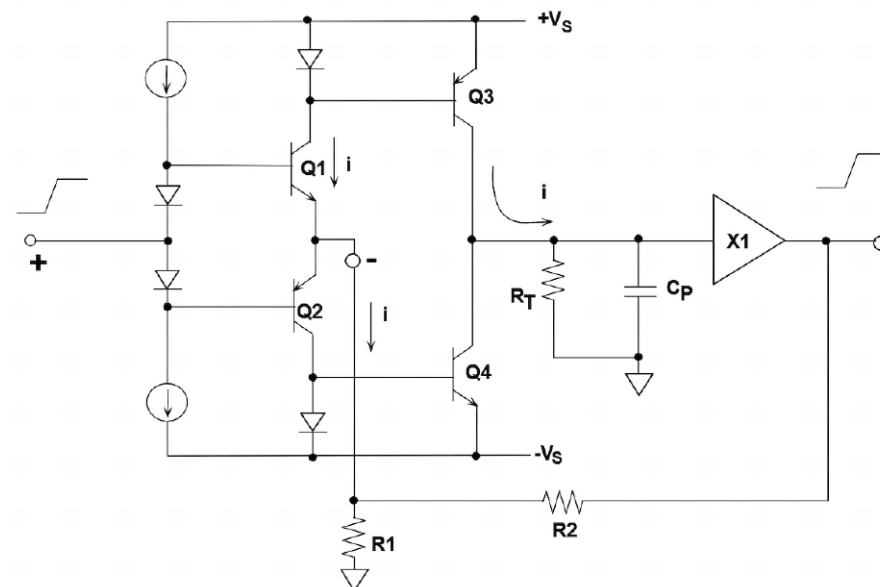
oraz wysokim slew-rate dla sygnałów o wysokich amplitudach. Taka szybkość pozwala na obsługę sygnałów w bardzo szerokim paśmie, co przybliża ów wzmacniacz do ideału pod względem szybkości działania (oczywiście, w zakresie do kilkunastu-kilkudziesięciu megaherców).

Grzechem byłoby nie wspomnieć o wzmacniaczach operacyjnych z prądowym sprzężeniem zwrotnym (CFB – current feedback), czyli takich, w którym wejście odwracające ma bardzo niską impedancję wejściową, rzędu pojedynczych omów. Uproszczony schemat takiego wzmacniacza wraz z rezystorami sprzężenia zwrotnego znajduje się na **rysunku 7**. Wynika z niego, że wejścia nie stanowi klasyczny układ różnicowy, lecz bazy komplementarnych tranzystorów bipolarnych (wejście nieodwracające, „+”) oraz ich emiteiry (wejście odwracające, „-”). Taki wzmacniacz operacyjny ma dużo więcej ograniczeń, jeżeli chodzi o jego otoczenie (rezystancje rezystorów, topologia), za to oferuje znacznie więcej niż typowy VFB (voltage feedback – wzmacniacz operacyjny z napiściowym sprzężeniem zwrotnym), przede wszystkim pod względem pasma i szybkości. Pierwszy z brzegu wzmacniacz operacyjny CFB, AD8011, oferuje pasmo sięgające ponad 100 MHz przy wzmacnieniu napiściowym +10 V/V.

Co słyszeć u komparatorów? Niewiele różni się pod tym względem od wzmacniaczy operacyjnych [9]. Mamy do dyspozycji zarówno różnicowe stopnie wejściowe z tranzystorami bipolarnymi (choćby popularny LM393 ze stopniem wejściowym niemal jak LM358 z **rysunku 2**), z tranzystorami



Rysunek 6. Konstrukcja stopnia wejściowego wzmacniacza rail-to-rail [6]



unipolarnymi JFET (bardzo fajny LF311, polecam) lub z tranzystorami unipolarnymi MOS (np. TLV3701 pobierający nieco ponad $0,5 \mu\text{A}$ prądu!). Warto nadmienić, iż nie istnieją komparatory CFB, ponieważ ta topologia służy linearyzacji pracy w szerokim zakresie częstotliwości, a przecież komparator jest „zwierzęciem” wybitnie nieliniowym.

Wzmacniacz napięciowy

Sama informacja z układu różnicowego to za mało, by powstał pełnoprawny wzmacniacz operacyjny. Ten słaby sygnał trzeba jeszcze wzmocnić w celu uzyskania podzespołu o tak wysokim (w teorii – nieskończenie wielkim) wzmocnieniu. Do realizacji wzmocnienia rzędu $60\ldots 100 \text{ dB}$ jest zaprzęgany z reguły jeden stopień w układzie wspólnego emitera lub wspólnego źródła. Całkiem czytelnie zostało to pokazane na schemacie wewnętrznym układu LM741 – rysunek 8.

Charakterystycznym fragmentem tego stopnia wzmacniającego (bądź stopni) jest obecność kondensatora o pojemności kilku pikofaradów lub większej (tu: $C_1 = 30 \text{ pF}$) między jego wejściem i wyjściem. Wejściem jest, w tym wypadku, baza tranzystora Q15, zaś wyjściem połączone kolektory tranzystorów Q15 i Q17. Obwód z tranzystorem NPN i rezystorami R7 i R8 służy jedynie przesunięciu potencjału, co umożliwia prawidłową polaryzację tranzystorów wyjściowych, lecz nie odgrywa istotnej roli w tej analizie. Obciążeniem jest wyjście lustra prądowego (kolektor tranzystora Q13) oraz bazy tranzystorów wyjściowych, Q14 i Q20. Zatem możliwe jest uzyskanie wysokiego wzmocnienia napięciowego z uwagi na wysoką rezystancję dynamiczną obciążenia „widzianego” przez kolektor.

Dodanie tego kondensatora powoduje zawężenie pasma, a to przez efekt Millera, który powoduje zwielokrotnienie wpływu jego pojemności (dzięki wysokiemu wzmocnieniu napięciowemu) na złącza baza-kolektor. Tak „uspokojony” układ staje się stabilny, przez wprowadzenie mu bieguna dominującego, zmniejszającego wzmocnienie do 1 V/V zanim przesunięcie fazowe wyniesie 180° – co doprowadziłoby do wzbudzenia układu objętego ujemnym sprzężeniem zwrotnym. Jest jednak wada: częstotliwość tego bieguna wynosi, na ogół, kilka-kilkanaście herców. Dla wspomnianego już 741, jest to około 5 Hz . Pięć herców! Stąd właśnie mój Promotor miał swoje „śmieszki” ze wzmacniaczy operacyjnych, jakoby były wolniejsze od zwykłych

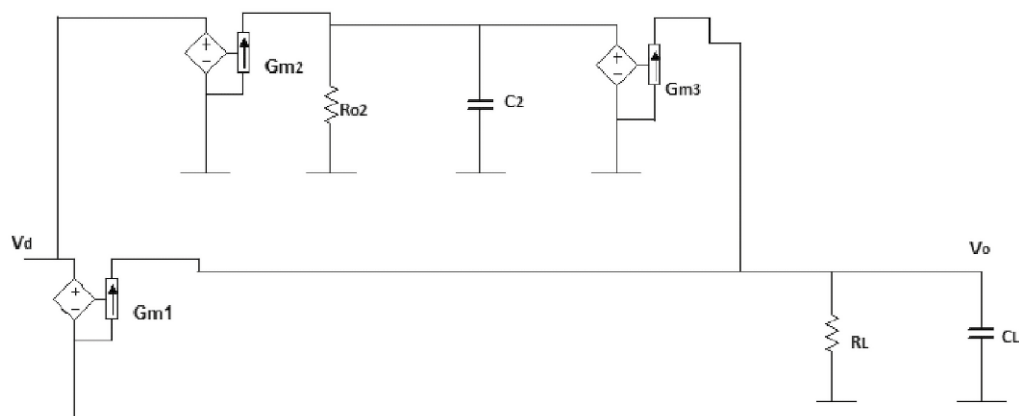
żarówek. Może i tak, ale w otwartej pętli sprzężenia zwrotnego używa się ich nader rzadko.

Wprowadzenie jednego bieguna dominującego jest bardzo wygodne: linearyzuje charakterystykę amplitudową, wymuszając jej spadek z nachyleniem 20 dB na dekadę, stabilizuje układ pracujący z dowolną pętlą ujemnego sprzężenia zwrotnego oraz upraszcza projektowanie dzięki wprowadzeniu zasady wymiany wzmocnienia na pasmo.

Nowoczesne wzmacniacze operacyjne są już, na szczęście, projektowane w inny sposób. Kompensacja biegunem dominującym była prosta w realizacji, ale relatywnie wolna jak na możliwości elementów półprzewodnikowych. Obecnie dominują układy wielobiegunowe (multi-pole lub pole-zero), która służy kształtowaniu charakterystyki przy zachowaniu stabilności. Jednak to nadal jest koncepcja polegająca na ograniczeniu pojedynczego wzmacniacza napięciowego o wysokim wzmocnieniu. Przykładem wzmacniacza operacyjnego o zaawansowanej kompensacji częstotliwościowej jest wspomniany już LT1364 o bardzo wysokim slew-rate [7].

Innym podejściem do rozwiązania tego problemu jest użycie topologii feedforward. Można ją znaleźć w wielu szybkich, nowoczesnych wzmacniaczach operacyjnych. Oprócz zwykłego, relatywnie wolnego toru przetwarzania sygnału, znajduje się w nim drugi, znacznie szybszy – rysunek 9. Zadaniem szybszego toru jest sterowanie wyjściem bezpośrednio ze stopnia wejściowego, co przyspiesza reakcję wyjścia, choć dzieje się to z mniejszym wzmocnieniem, niż ma to miejsce w przypadku wolniejszego toru o wysokim wzmocnieniu. Skutkuje to wyższym slew-rate oraz szerszym pasmem przenoszenia.

Dla kontrastu, wzmacniacze CFB oraz komparatory nie posiadają żadnego ogranicznika pasma w swoim stopniu napięciowym. Dla tych pierwszych nie jest to potrzebne, gdyż sama topologia zapewnia stabilność pod warunkiem prawidłowego doboru elementów, zaś komparatory celowo powinny być możliwie szybkie. W ich wypadku stosuje się stopnie napięciowe podobne do tych ze wzmacniaczy



Rysunek 9. Schemat blokowy dwustopniowego wzmacniacza operacyjnego z kompensacją feedforward [11]

operacyjnych, lecz bez dodatkowej kompensacji mającej na celu poprawę stabilności.

Stopień wyjściowy

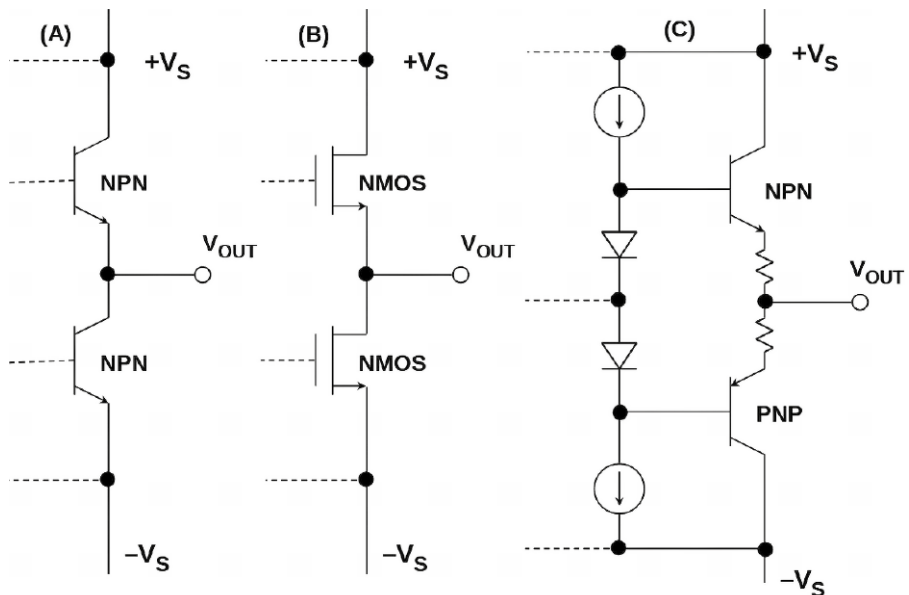
To trzeci najważniejszy „klocek” zarówno we wzmacniaczach operacyjnych, jak i w komparatorach. O dziwo, jest tutaj kilka wyraźnych różnic pomiędzy tymi typami układów scalonych. Ale – jak mawiał Bogusław Wołoszański – nie uprzedzajmy faktów.

Najpopularniejszym rodzajem stopnia wyjściowego jest push-pull zbudowany z dwóch komplementarnych tranzystorów: NPN+PNP w technologii bipolarnej lub NMOS+PMOS w technologii... a jakże, CMOS. NPN tudzież NMOS jest w stanie „dolewać” prąd do obciążenia z dodatniej linii zasilającej, natomiast PNP/PMOS jest w stanie go „wyciągać” w kierunku ujemnej linii zasilającej. Odpowiada to schematowi lewemu oraz środkowemu na rysunku 10. Dla zmniejszenia zniekształceń, stopnie te praktycznie zawsze pracują z niezerowym prądem spoczynkowym, czyli w klasie AB, co wymusza wstępne spolaryzowanie złącz baza-emiter (lub wymuszenie napięcia dren-źródło) obu tranzystorów, co może się odbywać choćby tak, jak na lewym schemacie rysunku 10: poprzez „rozsuniecie” potencjałów baz o napięcie zbliżone do dwukrotności napięcia przewodzącego złącza baza-emiter.

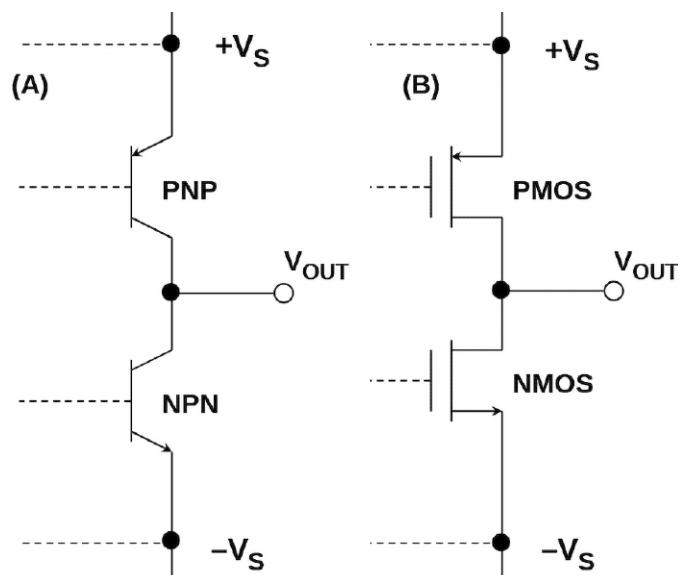
Co ciekawe, w tej topologii działa zdecydowana większość wzmacniaczy operacyjnych: od bardzo wolnych i przestarzałych, jak $\mu A741$ czy LM358, poprzez nowocześniejsze, jak OPA2134 czy wysokoprądowe jak OPA2544. Ta topologia zapewnia liniowość przy zachowaniu szybkości, zaś po obudowaniu dodatkowymi tranzystorami pełniącymi rolę wtórników może sterować „wymagającymi” obciążeniami. Z racji szybkości, taka topologia jest również stosowana we wzmacniaczach CFB.

Nieco inaczej, a wręcz: kompletnie na odwrót, jest wykonany stopień wyjściowy rail-to-rail, co udowadnia rysunek 11. Tranzystory nie pracują w nim w układzie wspólnego kolektora/drenu (czyli jako wtórniki napięciowe), lecz w konfiguracji wspólnego emitera/źródła. Pozwala to na „dociągnięcie” napięcia niemal do linii zasilającej, z dokładnością do napięcia nasycenia tranzystora bipolarnego lub spadku napięcia na rezystancji dren-źródło otwartego kanału tranzystora MOS.

Wiele osób zachwyca się tym typem wyjścia, ponieważ pozwala ono na osiągnięcie napięcia na wyjściu niemal równego potencjałowi którejś linii zasilającej,



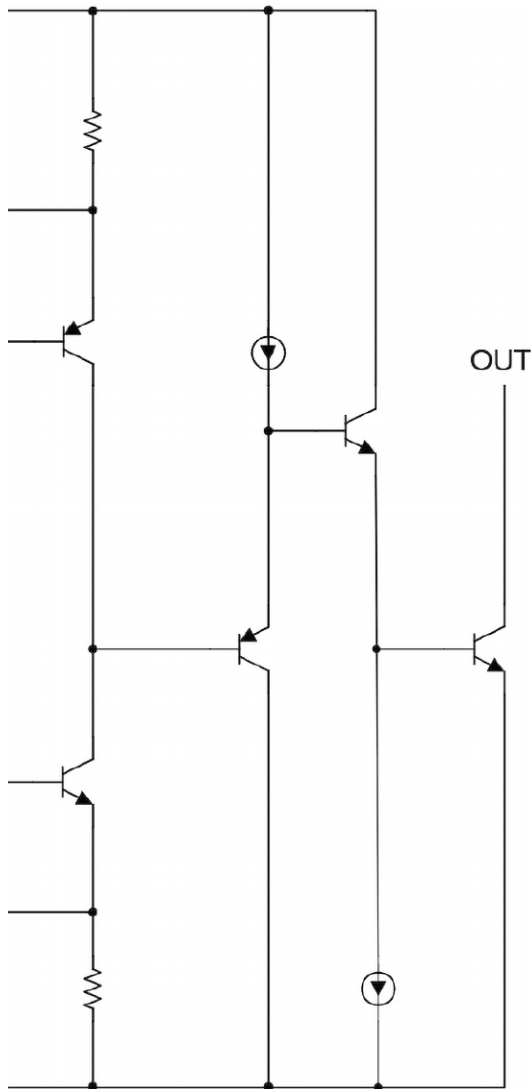
Rysunek 10. Schematy stopni wyjściowych w układzie push-pull [12]



Rysunek 11. Schematy stopni wyjściowych rail-to-rail [12]

z minimalnym odstępem. Przysłowiowe schody zaczynają się, kiedy wejdziemy w szczegóły: ów odstęp będzie tym większy, im większy będzie prąd wyjściowy – to dosyć oczywiste. Jeżeli tranzystor bipolarny wejdzie w nasycenie, osiągając ekstremum napięcia wyjściowego, wówczas dosyć wolno on z tego stanu wychodzi co ogranicza pasmo i wprowadza zniekształcenia. Dodać do tego należy również efekt Millera, który oddziałuje na aktywny w danej chwili tranzystor wyjściowy. Składając to wszystko razem, wolę uzyskać naddatek napięcia i użyć typowego wzmacniacza operacyjnego lub – jeżeli nie jest to możliwe, a taki efekt mnie zadowala – zastosować tani wzmacniacz CMOS.

Komparatory, czego pojąć do dzisiaj nie mogę, często mają wyjście typu otwarty kolektor (rysunek 12): zarówno proste i stare pokroju LM393, jak i nowocześniejsze, np. TLV1701. W wielu przypadkach, takie rozwiązanie jest mało wygodne: owszem, stan niski jest realizowany

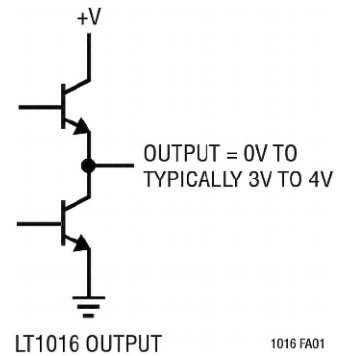


Rysunek 12. Wyjście typu open-collector układu TLV1701 [13]

szybko i z niezłą wydajnością prądową, za to wysoki... szkoda gadać. Utrudnia to wbudowanie histerezy, ponieważ zmniejszenie impedancji wyjściowej w stanie wysokim wymaga dodania wtórника, bądź też trzeba te skoki impedancji wyjściowej odpowiednio uwzględnić. W sporadycznych przypadkach taki rodzaj wyjścia się przydaje – przede wszystkim tam, gdzie napięcie sterowane przez komparator jest zupełnie inne niż to, które go zasilają.

Bardzo ciekawym rozwiązaniem, oferowanym przez producentów w szybkich (rzędu nanosekund) komparatorach, jest wyjście różnicowe, Q i Q̄. Ciekawym przykładem jest bardzo stabilny i dobrze opisany przez producenta LT1016 [14]. Ma ciekawy stopień wyjściowy, którego uproszczony schemat można zobaczyć na **rysunku 13**. Stan niski jest osiągnięty przez nasycający się „dolny” tranzystor NPN, zaś stan wysoki przez „górny” tranzystor NPN, pełniący z kolei rolę wtórника. W ten sposób układ jest w stanie uzyskać TTLowskie napięcia, czyli niemal 0 V w stanie niskim i 3...4 V w stanie wysokim. Może też sterować liniami różnicowymi, po zapewnieniu odpowiedniego dopasowania, bez konieczności dodawania wtórników.

Istnieją również proste, tanie i niezbyt szybkie komparatory z wyjściem push-pull, na przykład TLV3691 [15], pobierający jedynie 75 nA. Taki wynik, że proszę siadać! Dzięki takiemu wyjściu, może wymusić zmianę stanu logicznego np. na wejściu mikrokontrolera (w celu wybudzenia) bez angażowania w to rezystora podciągającego, który absorbowałby miejsce i dodatkowe, cenne mikroampery.



Rysunek 13. Stopień wyjściowy szybkiego komparatora LT1016 [14]

Układy nietypowe

Niektóre aplikacje wymagają obsługi wysokich napięć, spotkałem się z tym przede wszystkim w projektach sterowania przetworników piezoelektrycznych. Tam potrzebne są układy wysokonapięciowe. Przykładem bardzo ciekawego wzmacniacza, obsługującego aż do ± 50 V, jest OPA454. Jego stopień wyjściowy może obsługiwać prąd do 50 mA, przez co można nim sterować, bez dodatkowych wtórników, całkiem mocne tranzystory wyjściowe. Pomimo tego, oferuje całkiem szerokie pasmo i slew-rate pozwalający obsłużyć częstotliwości ponadakustyczne – szczegóły na **rysunku 14**.

Czymże byłby artykuł o wzmacniaczach operacyjnych bez choćby wspomnienia o wzmacniaczach instrumentalnych (instrumentation amplifiers – INA)? Nie wiem, więc jestem zobowiązany wspomnieć o bodaj najpopularniejszym – przynajmniej w moich układach – INA333. Względnie tani, względnie dobry, niezbyt szybki (przy $G = 100$ V/V jego pasmo wynosi zaledwie 3,5 kHz [17]), za to z wbudowanymi zabezpieczeniami RF na wejściach, przez co bardzo dobrze sprawdza się w aplikacjach przetwarzających sygnały wolnozmiennne, na przykład z czujników ciśnienia lub tensometrów. Zmiana wzmocnienia odbywa się, jak w każdym przystrojenym INA, za pomocą jednego rezystora – **rysunek 15**. Moim zdaniem, to jeden z najwygodniejszych układów do aplikacji, dodatkowo pobierający niewielki (niecałe 200 μ A) prąd zasilający.

Na sam koniec wspomnę jedynie o wzmacniaczach w pełni różnicowych (zarówno na wejściu, jak i na wyjściu) [18], które ułatwiają sterowanie linii różnicowych. Dzięki nim bardzo łatwo realizuje się konwertery sygnałów asymetrycznych (single-ended) na symetryczne (differential). Warto je znać, bo choć nie są tanie, to w zastosowaniach szybkich i wymagających niskich zniekształceń mogą pokazać, co potrafią.

PARAMETER	TEST CONDITIONS	MIN	TYP	MAX	UNIT
FREQUENCY RESPONSE⁽⁴⁾					
GBW Gain-bandwidth product	Small-signal		2.5		MHz
SR Slew rate	$G = \pm 1$, $V_O = 80$ -V step, $R_L = 3.27$ k Ω		13		V/ μ s
Full-power bandwidth ⁽⁵⁾			35		kHz
t_s Settling time ⁽⁶⁾	To $\pm 0.1\%$, $G = \pm 1$, $V_O = 20$ -V step		3		μ s
	To $\pm 0.01\%$, $G = \pm 5$ or ± 10 , $V_O = 80$ -V step		10		μ s
THD+N Total harmonic distortion + noise ⁽⁷⁾	$V_S = +40.6$ V/-39.6 V, $G = \pm 1$, $f = 1$ kHz, $V_O = 77.2$ V _{PP}		0.0008%		

Rysunek 14. Typowe parametry dynamiczne układu OPA454 [16]

Podsumowanie

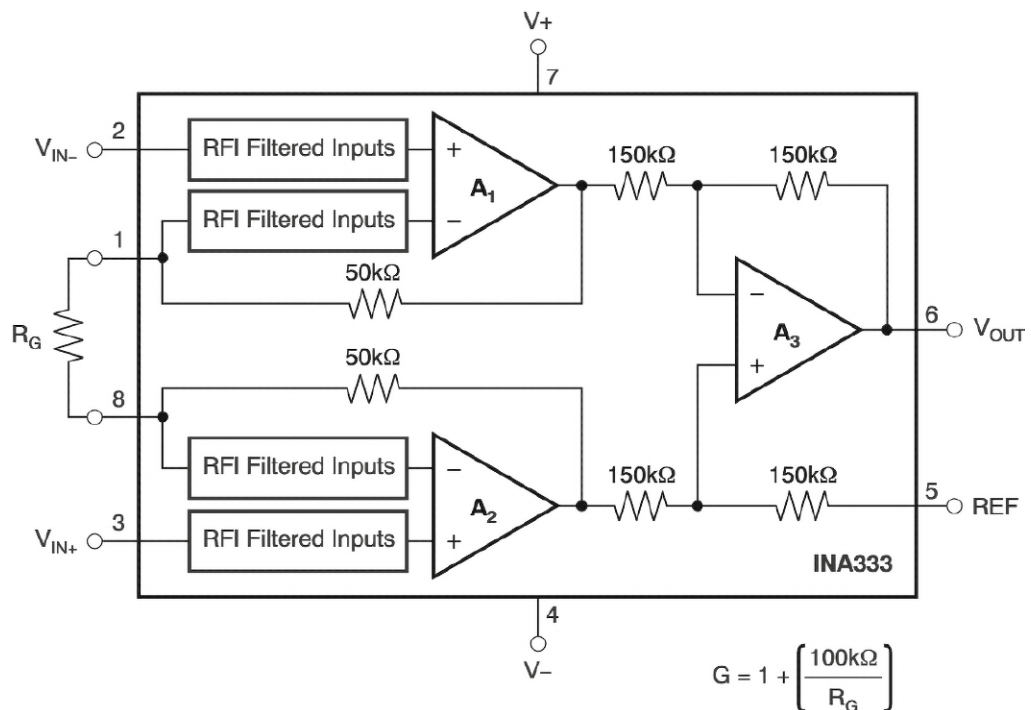
Wbrew opiniom niektórych, wzmacniacze operacyjne nie wybierają się na żadną emeryturę. Nadal chętnie stosowane, a dzięki osiągnięciom techniki półprzewodnikowej stają się coraz bardziej wszechstronne: niski pobór prądu, wysoka szybkość, niski prąd wejściowy to mo-
zownie „szlifowane” cechy, które zbliżają je do wzmacniaczy idealnych. Poza tym, mają również ciekawe rozwiązania, niegdyś niespotykane, jak wyjścia różnicowe czy bogaty wybór scalonych wzmacniaczy instrumentalnych, przydatnych choćby do pomiaru prądu czy przetwarzania sygnału z czujników.

Komparatory, będące spokrewnione ze wzmacniaczami operacyjnymi, również mają się dobrze i stanowią cenny „pomost” między światem analogowym a cyfrowym. Wszak bez nich nie mógłby powstać jakikolwiek przetwornik analogowo-cyfrowy! Konstruktorzy mają do wyboru bardzo szeroką ofertę ciekawych układów, w tym – co szczególnie cieszy moją osobę – z wyjściem push-pull, co znacznie upraszcza realizowanie histerezy i sterowanie układami cyfrowymi.

Michał Kurzela, EP

Źródła:

- [1] https://toshiba.semicon-storage.com/us/semiconductor/knowledge/faq/linear_opamp/what-is-the-ideal-op-amp.html
- [2] <https://www.onsemi.com/download/data-sheet/pdf/lm358-d.pdf>
- [3] <https://www.ti.com/lit/ds/symlink/tlo82-n.pdf>
- [4] <https://www.ti.com/lit/ds/symlink/opa336.pdf>



Rysunek 15. Schemat blokowy układu INA333 [17]

- [5] <https://www.ti.com/lit/ds/symlink/opa356-q1.pdf>
- [6] <https://www.ti.com/lit/an/sloa039a/sloa039a.pdf>
- [7] <https://www.analog.com/media/en/technical-documentation/data-sheets/13645fa.pdf>
- [8] <https://www.analog.com/media/en/training-seminars/tutorials/mt-057.pdf>
- [9] <https://www.ti.com/lit/an/snoaa91/snoaa91.pdf>
- [10] <https://www.ti.com/lit/ds/symlink/lm741.pdf>
- [11] <https://e-archivo.uc3m.es/rest/api/core/bitstreams/39066b1a-5d5e-4eeb-9bb0-883b9ed3beae/content>
- [12] <https://www.analog.com/media/en/training-seminars/tutorials/MT-035.pdf>
- [13] <https://www.ti.com/lit/ds/symlink/tlv1701.pdf?ts=1772613188429>
- [14] <https://www.analog.com/media/en/technical-documentation/data-sheets/lt1016.pdf>
- [15] <https://www.ti.com/lit/ds/symlink/tlv3691.pdf>
- [16] <https://www.ti.com/lit/ds/symlink/opa454.pdf>
- [17] <https://www.ti.com/lit/ds/symlink/ina333-ht.pdf>
- [18] <https://www.ti.com/lit/an/sloa054e/sloa054e.pdf>



Czujnik Shelly Flood Gen4 wykrywa wilgoć już w chwili pojawienia się wycieku i ostrzega, zanim dojdzie do poważniejszych szkód. Fot. Shelly

Shelly: ochrona przed zalaniem i zamek bez klucza



Na platformie zaopatrzeniowej Conrad pojawiły się dwie nowości marki Shelly – czujnik zalania Flood Gen4 oraz inteligentny zamek Shelly LOQED Touch Smart Lock. Oba urządzenia rozszerzają możliwości systemów Smart Home w biurze, warsztacie i w domu.

Shelly należy do czołowych dostawców rozwiązań dla inteligentnego domu. Urządzenia tej marki są wydajne, elastyczne i przystępne cenowo, a przy tym łatwe w obsłudze oraz instalacji – bez trudu można je włączyć do działających już systemów. Pracą można zarządzać przez aplikację, według harmonogramu lub głośnie. Sprzęt współpracuje z Alexa, Google Home i Home Assistant, można go zintegrować z ekosystemem Matter, a także korzystać z łączności Bluetooth, Wi-Fi i Zigbee.

Wczesne wykrywanie zalania

Nieszczelny wąż zmywarki, wyciek z instalacji wodociągowej czy ulewa zalewająca taras – czujnik Shelly Flood Gen4 w połączeniu z kablem Shelly Leak Sensor wykrywa takie zagrożenia już w chwili ich pojawienia się i wysyła sygnał alarmowy na smartfony oraz do systemów Smart Home. Dwumetrowy, elastyczny kabel czujnika można poprowadzić również w trudno dostępne miejsca; wystarczy ułożyć go na podłodze tam, gdzie ryzyko wycieku jest największe.

Automatyczne ograniczanie szkód

Flood Gen4 nie tylko powiadamia o zagrożeniu, ale w połączeniu z modułem Shelly 2PM Gen4 pozwala

zaprogramować scenariusze automatyki. W razie pęknięcia rury system może samodzielnie zamknąć główny zawór wody. Na zewnątrz czujnik rozpoznaje zbliżający się deszcz po rosnącej wilgotności powietrza – markiza składa się, zanim spadną pierwsze krople.

Adaptacyjne poziomy alarmu

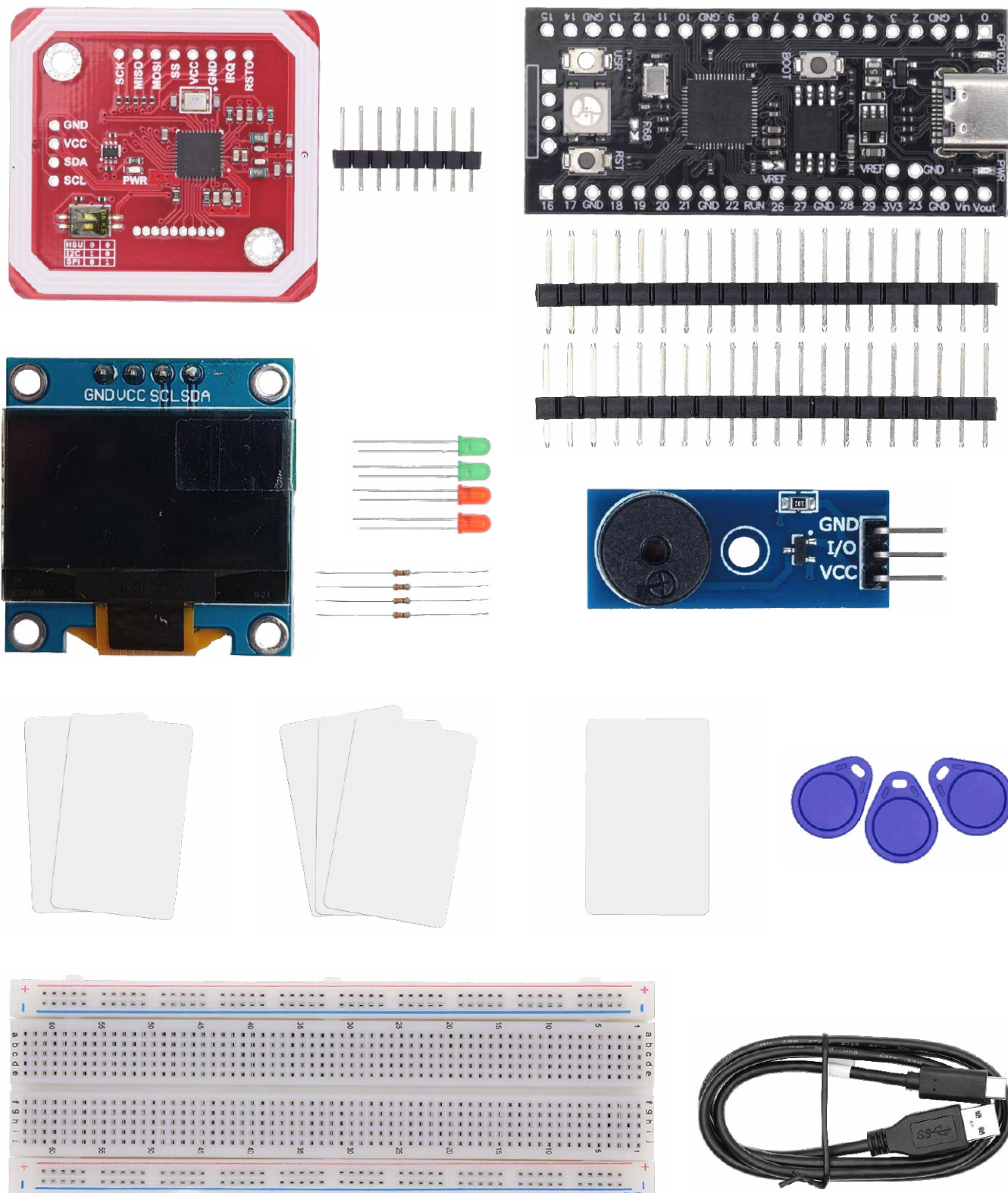
Czułość urządzenia można dopasować do specyfiki chronionego miejsca – serwerowni, domowej pralni czy rzadko używanego mieszkania. Dostępne są trzy tryby pracy: Intensywny, Normalny i Oszczędny. Czujnik pozwala też przełączać się między wykrywaniem zalania a wykrywaniem deszczu.

Zamek, który działa bez klucza

Shelly Loqed Touch Smart Lock umożliwia wejście za pomocą smartfona, kodu PIN lub lekkiego dotknięcia. Po połączeniu z ekosystemem Shelly zamek może być częścią scen automatyki – otwarcie drzwi uruchamia na przykład podniesienie rolet i włączenie światła. Zdalny dostęp pozwala wpuścić pracowników lub gości pod nieobecność domowników. Wykonany ze stali nierdzewnej, łatwy w montażu moduł sprawdza się w biurach, warsztatach oraz mieszkaniach.

Czujnik Shelly Flood Gen4 oraz zamek Shelly LOQED Touch Smart Lock są dostępne na platformie Conrad Sourcing Platform. Pełną ofertę marki Shelly można znaleźć na stronie conrad.pl.

Kurs RFID składa się z 10 rozdziałów. Rozdział 1, prezentujący bazę sprzętową kursu, publikujemy w tym wydaniu. Pozostałe 9 rozdziałów zamierzamy opublikować w najbliższych dwóch wydaniach. 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.



1. Zestaw sprzętowy do kursu RFID: moduł PN532, Raspberry Pi Pico, płytka stykowa, przewody, karty testowe. Źródło: fotografia własna redakcji/sklep.avt.pl

Kurs RFID z PN532, rozdział 1

RFID od zera: podłączamy PN532 i piszemy pierwszą klasę

Karty zbliżeniowe towarzyszą nam na każdym kroku. Karnet na siłownię, przepustka do biura, bilet miesięczny, a nawet klucz do hotelowego pokoju – to wszystko RFID. W tym kursie zajrzymy pod maskę tej technologii, zaczniemy od fizyki i skończymy na pisaniu interaktywnych wizytówek, które smartfon odczyta bez żadnej aplikacji. Zanim jednak dotkniemy czegokolwiek ciekawego, musimy porozmawiać z czytnikiem. I właśnie to jest temat dzisiejszego odcinka.

Co kupić przed pierwszym rozdziałem

Zanim podłączymy cokolwiek, sprawdź, czy masz komplet sprzętu. Kurs składa się z 10 rozdziałów – poniżej lista wszystkiego, czego będziesz potrzebować. Sprzęt z zestawu podstawowego wystarczy do rozdziałów 1...7. Elementy dodatkowe są potrzebne od rozdziału 8 wzwyż. Warto kupić wszystko od razu, bo to oszczędza czas i koszty wysyłki. Komplet wszystkich elementów potrzebnych do przerobienia całego kursu zawiera „Zestaw do kursu RFID” w sklepie AVT (sklep.avt.pl).

Zestaw podstawowy – rozdziały 1...7

Element	Uwagi
Moduł PN532 NFC RFID (Elechouse V3 lub V4)	Interfejsy SPI/I ² C/HSU, zasilanie 3,3...5 V, zasięg ok. 5 cm
Raspberry Pi Pico (wersja H – z przylutowanymi pinami)	Wersja H eliminuje konieczność lutowania pinów
Płytki stykowa (breadboard) 830 pól	Do połączeń bez lutowania
Przewody jumper – zestaw M-M i M-F, min. 20 szt.	Wymagane M-F do połączenia Pico z PN532
Kabel micro-USB do danych (nie tylko ładowania!)	Upewnij się że kabel przesyła dane – niektóre tylko ładują

Karty i tagi testowe

To najważniejsza część listy zakupów. Do różnych rozdziałów potrzebujesz różnych kart. Wszystkie znajdują się

wraz z kompletem pozostałych elementów w „Zestawie do kursu RFID” (sklep.avt.pl).

Karta/tag	Ilość i uwagi
MIFARE Classic 1k (oryginał NXP)	3 szt. – rozdziały 3, 4, 10
MIFARE Classic 1k Gen1a Magic Card (UID changeable)	3 szt. – rozdziały 4, 9, 10. Karta z możliwością zmiany UID – niezbędna w rozdziałach o bezpieczeństwie i emulacji.
MIFARE Classic 4k	1 szt. – rozdziały 3 i 4 (do porównania mapy pamięci)
MIFARE Ultralight (MF0ICU1)	2 szt. – rozdział 5
MIFARE Ultralight EV1	2 szt. – rozdział 5 (do porównania z Ultralight)
NTAG216	3 szt. – rozdziały 6, 7, 10
MIFARE DESFire EV1 lub EV2	1 szt. – rozdział 8. Można zastąpić własną kartą miejską lub bankomatową (SAK=0x20).

Do projektu końcowego (rozdział 10) i demonstracji emulacji kart (rozdział 9) potrzebujesz kilku niżej wymienionych elementów elektronicznych dostępnych wraz z pozostałymi elementami w „Zestawie do kursu RFID” (sklep.avt.pl).

Element	Uwagi
Wyświetlacz OLED 0,96" I ² C, 128×64 px, sterownik SSD1306	Wersja z 4 pinami (GND/VCC/SCL/SDA) – najprostsza do podłączenia przez I ² C
Dioda LED zielona + czerwona (5 mm)	Po 2 sztuki
Rezystor 330 Ω (do LED)	4 sztuki
Aktywny buzzer (moduł 3,3 V lub 5 V)	Moduł wygodniejszy niż goły buzzer – nie wymaga dodatkowych elementów
Dodatkowe przewody jumper M-F (ok. 20 szt.)	OLED i projekt końcowy zajmują wiele pinów jednocześnie

Smartfon z NFC (Near Field Communication) zastępuje drugi czytnik

Do testowania emulacji kart (rozdział 9) i odczytywania wizytówek NDEF (rozdział 7) wystarczy smartfon z NFC – Android od wersji 4.0, iPhone od XS. Sprawdź w ustawieniach telefonu, czy masz NFC i czy jest włączone. Nie masz telefonu z NFC? Jedynym wyjściem jest drugi moduł PN532 jako czytnik. Kup go jako element dodatkowy razem z „Zestawem do kursu RFID”, zamawiając „Zestaw do kursu RFID plus” (sklep.avt.pl).

Cecha	LF (Low Freq.)	HF (High Freq.)	UHF (Ultra High Freq.)
Częstotliwość	125...135 kHz	13,56 MHz	860...960 MHz
Zasięg	do 10 cm	do 1 m (typowo 5...20 cm)	1...10 m i więcej
Standard	brak dominującego	ISO 14443, ISO 15693	ISO 18000-6, EPC Gen2
Typowe karty	EM4100, HID Prox	MIFARE, DESFire, NTAG	etykiety logistyczne
Zastosowania	dostęp, zwierzęta	płatności, bilety, NFC	magazyny, paczki
Prędkość danych	wolna	ok. 106...424 kb/s	szybka

Trzy rodziny RFID

Skrót RFID pochodzi od *Radio Frequency Identification* – identyfikacja za pomocą fal radiowych. Technologia istnieje od lat 70. ubiegłego wieku, ale do powszechnego użytku trafiła dopiero dwie dekady później. Dziś dzieli się na trzy główne rodziny, różniące się częstotliwością pracy, zasięgiem i typowymi zastosowaniami.

Nas interesuje środkowa kolumna. Częstotliwość 13,56 MHz to świat kart MIFARE, NTAG, DESFire i ogólnie NFC. Tu obowiązuje norma ISO 14443 – i to na jej bazie zbudowany jest cały ten kurs.

Jak karta działa bez baterii

Jedno z najczęstszych pytań, jakie słyszę od osób stawiających pierwsze kroki z RFID, brzmi mniej więcej tak: „jak karta w ogóle może coś wysłać, skoro nie ma baterii?”

Odpowiedź kryje się w zjawisku indukcji elektromagnetycznej odkrytym przez Faradaya w 1831 roku. Czytnik wytwarza zmienne pole magnetyczne o częstotliwości 13,56 MHz. Antena karty – to zwykła płaska cewka – znajduje się w tym polu i na jej końcach indukuje się napięcie. Po prostowaniu i stabilizacji zasila ono układ scalony karty. Karta budzi się, inicjalizuje pamięć i jest gotowa na polecenia – wszystko w czasie krótszym niż milisekunda.

Transmisja danych odbywa się w obie strony, ale zupełnie różnymi metodami:

Jak duże jest napięcie indukowane?

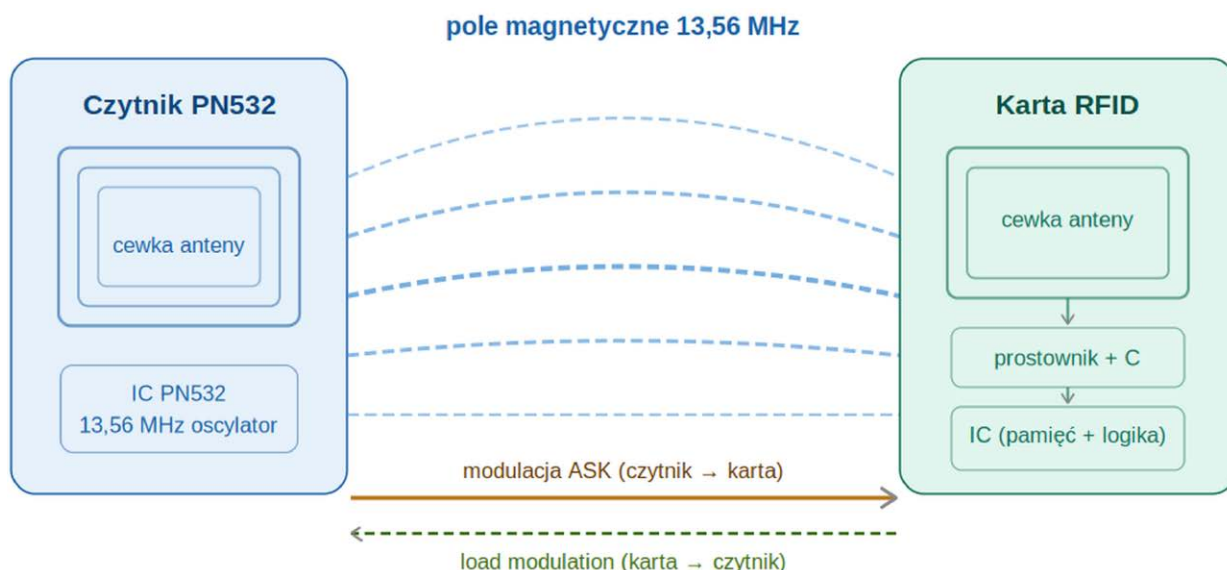
Cewka anteny w karcie ISO 14443A ma typowo 4...6 zwojów i kondensator dobrany dla rezonansu, który sprawia, że układ LC rezonuje dokładnie przy 13,56 MHz. W rezonansie napięcie na końcach cewki może sięgać kilku woltów, mimo że pole czytelnika jest dość słabe. Prostownik i stabilizator napięcia – zintegrowane w układzie karty – dostarczają zasilanie na poziomie około 1,8...3,3 V dla logiki.

- **Czytnik → karta:** modulacja ASK (*Amplitude Shift Keying*). Czytnik na chwilę „przydusza” amplitudę pola o 10% lub 100% – to koduje 0 i 1. Karta widzi zmianę napięcia na swojej cewce i interpretuje ją jako bit danych.
- **Karta → czytnik:** load modulation. Karta przyłącza mały rezystor równolegle do swojej anteny. Zmiana obciążenia pola jest mała, ale czuły układ czytelnika ją wykrywa. Tak karta „odpowiada” bez żadnego własnego nadajnika.

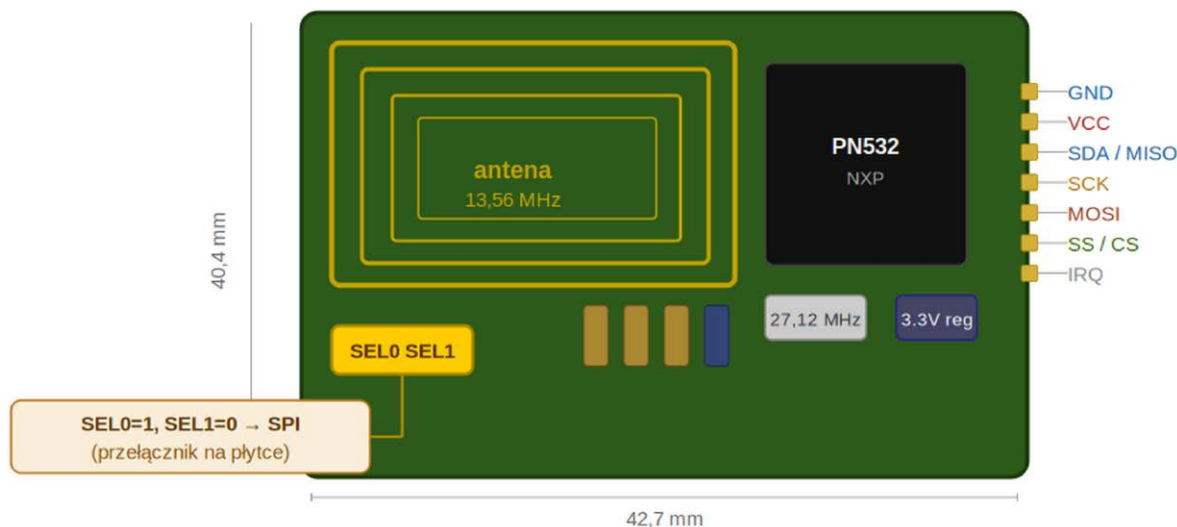
Czytnik PN532 – dlaczego akurat ten?

Do tego kursu wybrałem moduł oparty na układzie scalonym PN532 firmy NXP. Dlaczego nie popularny i tani RC522? Kilka powodów.

RC522 jest układem przestarzałym. Nadal spełnia swoje zadanie, ale ma sporo ograniczeń: obsługuje tylko tryb reader/writer (nie potrafi emulować karty ani robić peer-to-peer), pracuje wyłącznie przez SPI i nie ma wbudowanej



2. Zasada działania RFID 13,56 MHz: pole magnetyczne czytnika zasila kartę indukcyjnie; modulacja ASK (czytnik→karta) i load modulation (karta→czytnik). Źródło: opracowanie własne na podstawie: Wikimedia Commons, CC0 (commons.wikimedia.org/wiki/File:Wireless_power_system_-_inductive_coupling.svg)



3. Schemat modułu Elechouse PN532 NFC RFID Module V3: antena spiralna, układ PN532, przełączniki SEL0/SEL1, goldpin.
Wymiary: 42,7 × 40,4 mm. Źródło: opracowanie własne na podstawie: Elechouse PN532 NFC RFID Module V3 User Guide (manuals.plus/elechouse/pn532-nfc-rfid-module-manual)

obsługi NDEF. PN532 jest nowocześniejszy – obsługuje ISO 14443A/B, ISO 15693, tryb emulacji karty, komunikację peer-to-peer (czyli to, co w telefonie nazywamy NFC) i dostępny jest przez SPI, I²C lub UART. Kosztuje trochę więcej niż RC522, ale moduł z nim jest bez problemu dostępny w polskich sklepach za kilkanaście złotych. Jest oczywiście podstawowym elementem w „Zestawie do kursu RFID” (sklep.avt.pl).

Cecha	PN532
Producent	NXP Semiconductors
Rok wprowadzenia	2004 (aktywnie rozwijany)
Standardy	ISO 14443A/B, ISO 15693, NFC IP-1
Tryby pracy	reader/writer, card emulation, P2P
Interfejsy	SPI, I ² C, HSU (UART)
Napięcie zasilania	3,3 V lub 5 V (przez wbudowany regulator)
Zasięg	typowo 5 cm, do 10 cm w korzystnych warunkach
Cena modułu	w „Zestawie do kursu RFID” (sklep.avt.pl)

Popularny moduł to Elechouse PN532 NFC RFID Module V3 – płytka zawierająca PN532, antenę i stabilizator napięcia. Na module są dwa małe przełączniki (*jumper*

Ustawienie przełączników SEL0 / SEL1 na module PN532

SEL0	SEL1	Tryb komunikacji
1 (ON)	0 (OFF)	SPI – kurs używa tego
0 (OFF)	0 (OFF)	HSU (UART)
0 (OFF)	1 (ON)	I ² C

Przełączniki SMD na krawędzi płytki — przesunąć w pozycję oznaczoną na PCB

4. Ustawienie przełączników SEL0/SEL1 dla każdego trybu komunikacji modułu PN532. Kurs używa trybu SPI: SEL0 = 1 (ON), SEL1 = 0 (OFF). Źródło: opracowanie własne na podstawie: Elechouse PN532 User Guide i dokumentacji NXP PN532 (nxp.com – PN532 User Manual UM10232)

Gdzie kupić zestaw?

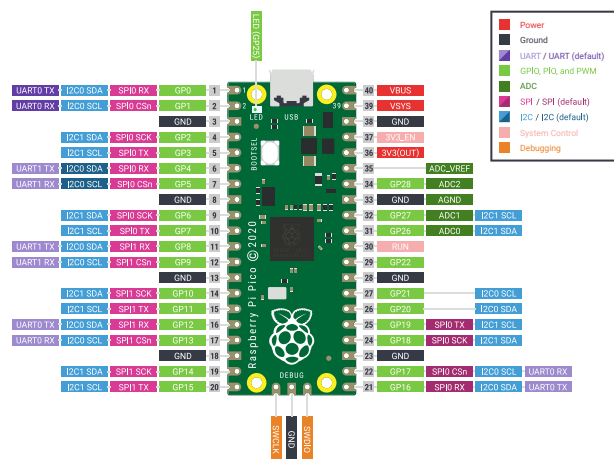
Po weryfikacji kursu kompletny „Zestaw do kursu RFID” będzie dostępny w sklepie.avt.pl.

switch), którymi wybieramy interfejs komunikacyjny. W tym kursie korzystamy z SPI (*Serial Peripheral Interface*).

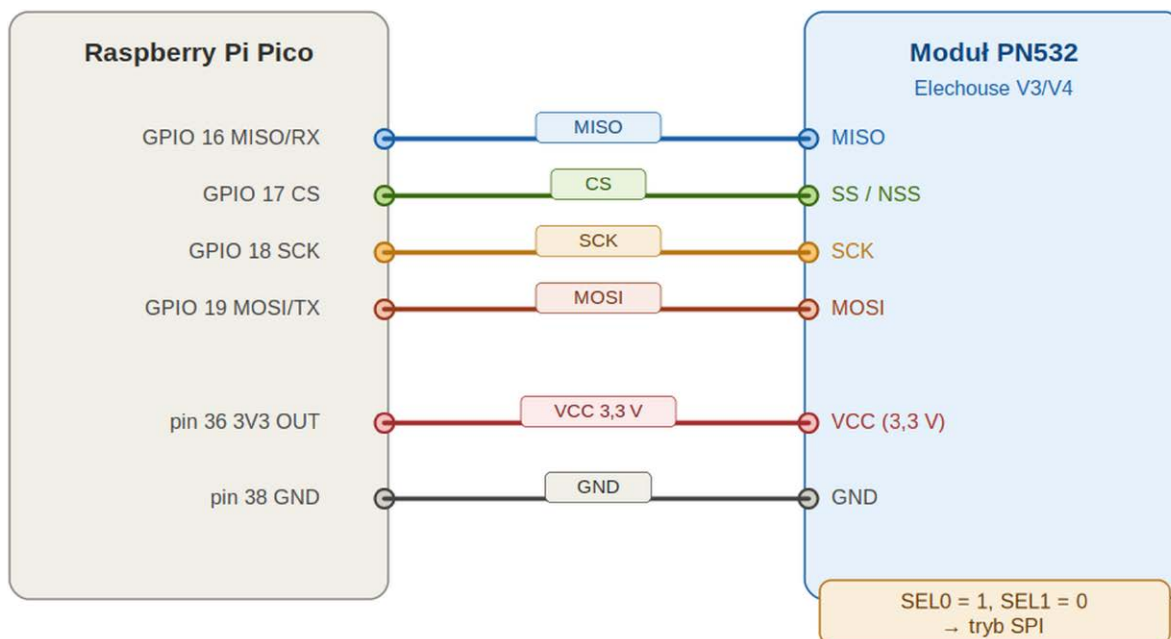
Podłączenie PN532 do Raspberry Pi Pico Ustawienie trybu SPI

Zanim podłączymy cokolwiek, musimy ustawić tryb komunikacji. Na module Elechouse PN532 V3 znajdziesz dwa małe przełączniki oznaczone SEL0 i SEL1 (lub SW1/SW2, zależnie od wersji). Ich kombinacja decyduje o interfejsie. Ustaw je zgodnie z tabelką na rysunku 4.

Jeśli moduł ma jumper zamiast przełącznika, zwórkę umieścić na pozycji SPI. Szczegółowy opis znajdziesz w instrukcji modułu (*Elechouse PN532 User Guide*).



5. Raspberry Pi Pico – piny używane w kursie: GPIO 16 MISO, GPIO 17 CS, GPIO 18 SCK, GPIO 19 MOSI (SPI0), 3V3 OUT (pin 36), GND (pin 38). Źródło: opracowanie własne na podstawie oficjalnego pinu Raspberry Pi Ltd (CC BY-ND 4.0) (datasheets.raspberrypi.com/pico/Pico-R3-A4-Pinout.pdf)



6. Schemat połączeń Raspberry Pi Pico ↔ PN532 przez SPI. Kolory: niebieski = MISO, zielony = CS, żółty = SCK, pomarańczowy = MOSI, czerwony = VCC 3,3 V, czarny = GND. Źródło: opracowanie własne; schemat poglądowy wzorowany na: Adafruit Industries, CC BY-SA (learn.adafruit.com/raspberry-pi-nfc-minecraft-blocks/hardware-wiring)

Schemat połączeń

Raspberry Pi Pico ma dwa kontrolery SPI: SPI0 i SPI1. Używamy SPI0 z domyślnymi pinami. Poniżej pełna tabela połączeń:

Sygnat SPI	Numer GPIO (Pico)	Numer pinu (Pico)	Pin na module PN532
SCK (zegar)	GPIO 18	pin 24	SCK
MOSI	GPIO 19	pin 25	MOSI
MISO	GPIO 16	pin 21	MISO
CS	GPIO 17	pin 22	NSS/SS
VCC	–	pin 36 (3V3)	VCC (3,3 V)
GND	–	pin 38 (GND)	GND

UWAGA: napięcie zasilania

Moduł PN532 akceptuje zasilanie 3,3 V lub 5 V (ma wbudowany regulator). Raspberry Pi Pico ma jednak linie sygnałowe 3,3 V. Podłącz VCC modułu do pinu 3V3 Pico (pin 36), nie do VBUS (5 V, pin 40) – unikniesz ryzyka przekroczenia poziomów logicznych na liniach SPI.

Jeśli używasz ESP32, możesz zasilac z 3,3 V lub 5 V – ESP32 toleruje 3,3 V na pinach GPIO, więc i tu wybierz 3V3.

Połączenia robimy na płytce stykowej (breadboard). Nie potrzeba lutowania. Wystarczy 6 przewodów jumper typu M-M lub M-F, w zależności od tego, czy masz Pico z przylutowanymi pinami czy na gołym module.

Weryfikacja połączeń przed uruchomieniem kodu

Zanim wgrasz cokolwiek, warto sprawdzić połączenia multimetrem oraz wzrokowo:

- Czy VCC na module ma 3,3 V względem GND? (możesz zmierzyć po podłączeniu USB do Pico)
- Czy CS podłączony jest do GPIO 17, a nie do 3V3 lub GND?
- Czy MOSI i MISO nie są zamienione miejscami? (to najczęstszy błąd)
- Czy jumper/przełącznik na module jest ustawiony na SPI?

Jak PN532 rozmawia przez SPI

Zanim napiszemy kod, musimy zrozumieć, jak wygląda komunikacja między Pico a PN532. Nie jest to zwykły SPI, gdzie wysyłasz bajty i od razu dostajesz odpowiedź. PN532 ma własny protokół warstwy aplikacyjnej, zbudowany nad SPI.

Format ramki PN532

Każda komenda i każda odpowiedź to „ramka normalna” (Normal Frame) opisana w PN532 User Manual, rozdz. 6.2.1. Jej budowa wygląda tak:

Pole	Wartość/opis
Preamble	0x00 – jeden bajt zerowy na początek
Start Code	0x00, 0xFF – dwa bajty znacznika początku
LEN	liczba bajtów payload (TFI + dane)
LCS	dopełnienie LEN do 256: $(LEN + LCS) \bmod 256 = 0$
TFI	0xD4 = host → PN532, 0xD5 = PN532 → host
CMD	kod komendy (np. 0x02 dla GetFirmwareVersion)
DATA	dane komendy (opcjonalne, różna długość)
DCS	suma kontrolna danych: $(TFI + CMD + DATA + DCS) \bmod 256 = 0$
Postamble	0x00 – jeden bajt zerowy na koniec

Dla przykładu: komenda GetFirmwareVersion (0x02) nie ma żadnych argumentów. Jej ramka wygląda tak:

```

Bajt:  00  00  FF  02  FE  D4  02  2A  00
      PRE SC1 SC2 LEN LCS  TFI  CMD  DCS  POST

```

Obliczenia sum kontrolnych:

LCS: $(0x100 - \text{LEN}) \& 0xFF = (0x100 - 0x02) \& 0xFF = 0xFE$

DCS: $(0x100 - (0xD4 + 0x02)) \& 0xFF = (0x100 - 0xD6) \& 0xFF = 0x2A$

Dlaczego SPI w PN532 jest LSB first?

PN532 używa SPI w niekonwencjonalny sposób: bajty są wysyłane od najmłodszego bitu (LSB first), podczas gdy standardowe SPI wysyła od najstarszego (MSB first). Na szczęście moduł Elechouse V3 ma na płycie układ odwracający kolejność bitów, więc z poziomu kodu MicroPython możesz używać domyślnego MSB. Jeśli jednak używasz surowego układu PN532 lub innego modułu, sprawdź w jego dokumentacji, czy potrzebujesz parametru `firstbit=SPI.LSB`.

Cykl zapis-ACK-odpowiedź

Protokół PN532 na SPI przebiega w trzech krokach. Błąd w jednym z nich objawia się jako brak odpowiedzi lub błędne dane:

- **Krok 1 – ZAPIS:** Host (Pico) obniża CS, wysyła bajt `0x01` (`SPI_DATAWRITE`) a po nim pełną ramkę komendy, po czym zwalnia CS.
- **Krok 2 – ACK:** PN532 przetwarza komendę i wystawia status `READY`. Host sprawdza go, wysyłając bajt `0x02` (`SPI_STATREAD`). Gdy `READY=1`, host odczytuje 6-bajtowy ACK: `00 00 FF 00 FF 00`. To potwierdzenie, że komenda dotarła poprawnie.
- **Krok 3 – ODPOWIEDŹ:** PN532 przetwarza komendę (może to trwać kilkadziesiąt ms) i ponownie wystawia `READY`. Host czyta odpowiedź, wysyłając bajt `0x03`

(`SPI_DATAREAD`) i tyle bajtów–zaślepek `0x00`, ile spożewia się danych.

Piszemy klasę PN532

Czas na kod. Będziemy tworzyć dwa pliki: `pn532.py` z klasą warstwy 1 i `main.py` z programem testowym. Oba wgrasz do pamięci Pico przez Thonny (File → Save As → Raspberry Pi Pico).

Zacznijmy od `pn532.py`. Czytaj kod razem z komentarzami – każdy fragment jest wyjaśniony. Wrócimy do tego pliku w kolejnych rozdziałach, dodając kolejne komendy.

Listing 1: `pn532.py` – klasa PN532 (warstwa 1)

```

# pn532.py
# Warstwa 1: komunikacja z układem PN532 przez SPI
# Kurs RFID, rozdział 1

from machine import Pin, SPI
import time

# -----
# Stałe protokołu PN532 (źródło: PN532 User Manual, UM10232, rozdz. 6.2.1)
# -----

_PREAMBLE      = 0x00 # bajt preambuły
_STARTCODE1    = 0x00 # pierwszy bajt kodu startowego
_STARTCODE2    = 0xFF # drugi bajt kodu startowego
_POSTAMBLE     = 0x00 # bajt zakończenia ramki

_HOST_TO_PN532 = 0xD4 # TFI: ramka idzie od nas DO PN532
_PN532_TO_HOST = 0xD5 # TFI: ramka idzie OD PN532 do nas

# Nagłówki transakcji SPI
_SPI_STATREAD  = 0x02 # zapytaj o gotowość danych
_SPI_DATAWRITE = 0x01 # wyślij dane do PN532
_SPI_DATAREAD  = 0x03 # odbierz dane od PN532

CMD_GET_FIRMWARE_VERSION = 0x02

class PN532:
    """Warstwa 1: niskopoziomowa komunikacja z PN532 przez SPI."""

    def __init__(self, spi, cs_pin):
        self._spi = spi
        self._cs = Pin(cs_pin, Pin.OUT, value=1) # CS nieaktywny = wysoki
        time.sleep_ms(100) # daj PN532 czas na uruchomienie

    # -----
    # Warstwa SPI: surowy zapis i odczyt bajtów
    # -----

```



```

def _cs_on(self):
    """Aktywuj CS (niski stan = PN532 zaznaczony jako rozmówca)."""
    self._cs(0)
    time.sleep_us(10) # PN532 potrzebuje chwili po CS

def _cs_off(self):
    """Zwolnij CS."""
    time.sleep_us(10)
    self._cs(1)

def _spi_write(self, data):
    """Wyślij listę bajtów do PN532."""
    self._cs_on()
    self._spi.write(bytes(data))
    self._cs_off()

def _spi_read_status(self):
    """Sprawdź, czy PN532 ma gotową odpowiedź.
    Zwraca True, jeśli TAK."""
    self._cs_on()
    buf = bytearray(2)
    self._spi.write_readinto(bytes([_SPI_STATREAD, 0x00]), buf)
    self._cs_off()
    return buf[1] == 0x01

def _spi_read(self, length):
    """Odbierz 'length' bajtów danych od PN532."""
    self._cs_on()
    tx = bytes([_SPI_DATAREAD] + [0x00] * length)
    rx = bytearray(len(tx))
    self._spi.write_readinto(tx, rx)
    self._cs_off()
    return rx[1:] # pierwszy bajt to echo nagłówka, pomijamy

# -----
# Warstwa ramek: budowanie i parsowanie protokołu PN532
# -----

def _build_frame(self, cmd, data=b''):
    """Zbuduj normalną ramkę PN532.
    Struktura (UM10232, rozdz. 6.2.1):
    00 00 FF [LEN] [LCS] [TFI] [CMD] [DATA...] [DCS] 00
    """
    payload = bytes([_HOST_TO_PN532, cmd]) + bytes(data)
    length = len(payload)
    lcs = (~length + 1) & 0xFF # suma LEN+LCS musi wynosić 0x00
    dcs = (~sum(payload) + 1) & 0xFF # suma bajtów payload+DCS = 0x00
    return bytes([
        _PREAMBLE, _STARTCODE1, _STARTCODE2,
        length, lcs,
    ]) + payload + bytes([dcs, _POSTAMBLE])

def _wait_ready(self, timeout_ms=1000):
    """Poczekaj, aż PN532 skończy przetwarzać i będzie gotowy."""
    t0 = time.time()
    while not self._spi_read_status():
        if time.time() - t0 > timeout_ms:
            raise TimeoutError('PN532 nie odpowiada - sprawdź połączenie')
        time.sleep_ms(10)

def _send_command(self, cmd, data=b''):
    """Wyślij komendę i poczekaj na ACK od PN532."""
    frame = self._build_frame(cmd, data)
    self._spi.write([_SPI_DATAWRITE] + list(frame))
    self._wait_ready()
    ack = self._spi_read(6)
    if bytes(ack) != bytes([0x00, 0x00, 0xFF, 0x00, 0xFF, 0x00]):
        raise RuntimeError(f'Zły ACK: {list(ack)}')

def _read_response(self, cmd, data_length):
    """Odbierz odpowiedź na komendę cmd. Zwraca bajty danych."""
    self._wait_ready()
    raw = self._spi_read(7 + data_length + 2)
    tfi = raw[4]
    rcmd = raw[5]
    if tfi != _PN532_TO_HOST or rcmd != (cmd + 1):
        raise RuntimeError(f'Nieoczekiwana odpowiedź: TFI={tfi:#x}, CMD={rcmd:#x}')
    return raw[6:6 + data_length]

```

```
# -----
# Komendy publiczne
# -----

def get_firmware_version(self):
    """Pobierz wersję firmware PN532.
    Zwraca krotkę (IC, Ver, Rev, Support).
    Przykład: (5, 1, 6, 7) oznacza PN532, wersja 1.6."""
    self._send_command(CMD_GET_FIRMWARE_VERSION)
    data = self._read_response(CMD_GET_FIRMWARE_VERSION, data_length=4)
    return tuple(data)
```

Kilka ważnych uwag do listingu:

- **Metody z prefiksem _ (podkreślnik)** to metody prywatne. Nie będziesz ich wywoływał z zewnątrz – to wewnętrzna maszyna klasy. Publiczne metody (jak `get_firmware_version`) to właściwy interfejs.
- **_build_frame** buduje ramkę dokładnie według specyfikacji z UM10232 rozdz. 6.2.1. Możesz porównać ją z tabelą powyżej – każda linia kodu odpowiada jednemu polu ramki.

- **wait_ready** odpytuje PN532 w pętli co 10 ms (polling). W przyszłości można by podpiąć pin IRQ modułu do przerywania sprzętowego, ale dla prostoty kursu polling wystarczy.
- **TimeoutError** zostanie podany, jeśli PN532 nie odpowie w ciągu 1 sekundy. To zwykle znaczy: błędne połączenie, zły tryb (I²C zamiast SPI) albo problem z zasilaniem.

Listing 2: main.py – test połączenia

```
# main.py
# Test połączenia z PN532 – uruchom jako pierwszy program

from machine import SPI, Pin
from pn532 import PN532

# -----
# Połączenia SPI0 na Raspberry Pi Pico:
# SCK (zegar)      -> GPIO 18 -> pin 24 Pico
# MOSI (dane do PN) -> GPIO 19 -> pin 25 Pico
# MISO (dane od PN) -> GPIO 16 -> pin 21 Pico
# CS (wybór układu) -> GPIO 17 -> pin 22 Pico
# -----

spi = SPI(0,
          baudrate=1_000_000, # 1 MHz – bezpieczna wartość na start
          polarity=0,         # CPOL=0: zegar w stanie niskim gdy bezczynny
          phase=0,            # CPHA=0: próbkowanie na zboczu narastającym
          sck=Pin(18),
          mosi=Pin(19),
          miso=Pin(16))

czytnik = PN532(spi=spi, cs_pin=17)

print('Łączenie z PN532...')
try:
    ic, ver, rev, support = czytnik.get_firmware_version()
    print(f' Układ:      PN5{ic:02X}')
    print(f' Wersja:    {ver}.{rev}')
    print(f' Support:   0b{support:08b}')
    print('OK – PN532 odpowiada poprawnie!')
except TimeoutError as e:
    print(f'Błąd: {e}')
    print('Sprawdź przewody i ustawienie zworek.')
except RuntimeError as e:
    print(f'Błąd protokołu: {e}')
```

Pierwsze uruchomienie Wgrywanie plików

Otwórz Thonny. Upewnij się, że interpreter jest ustawiony na MicroPython (Raspberry Pi Pico) – widać to w prawym dolnym rogu okna. Jeśli nie, przejdź do Run → Select Interpreter i wybierz MicroPython (Raspberry Pi Pico).

- Utwórz nowy plik (Ctrl+N), wklej zawartość **Listing 1** i zapisz jako `pn532.py` na Pico (File → Save As → Raspberry Pi Pico → wpisz nazwę `pn532.py`).
- Utwórz drugi plik z **Listingiem 2** i zapisz jako `main.py` na Pico.
- Uruchom `main.py` klawiszem F5 (lub zielonym trójkątem Run).

Oczekiwany wynik

W konsoli Thonny powinna pojawić się informacja podobna do poniższej (dokładne wartości zależą od wersji firmware modułu):

```
Łączenie z PN532...
Układ: PN532
Wersja: 1.6
Support: 0b00000111
OK - PN532 odpowiada poprawnie!
```

Wartość Support to bajt bitowy: bit 0 = ISO 14443A, bit 1 = ISO 14443B, bit 2 = ISO 18693. Wynik 0b00000111 oznacza obsługę wszystkich trzech – to standardowa wartość dla modułów Elechouse.

Co zrobić, gdy nie działa?

Jeśli zamiast powyższego zobaczysz błąd, sprawdź po kolei:

Tip: podgląd plików w Thonny

Thonny ma widok systemu plików Pico (View → Files). Po lewej widzisz pliki na komputerze, po prawej – pliki na Pico. To łatwy sposób sprawdzenia, czy oba pliki (pn532.py i main.py) są faktycznie na Pico, a nie tylko otwarte lokalnie.

Komunikat/objaw	Prawdopodobna przyczyna i rozwiązanie
TimeoutError: PN532 nie odpowiada	1. Sprawdź, czy CS jest na GPIO 17, nie na 3V3. 2. Sprawdź ustawienie przetącnika na SPI. 3. Sprawdź zasilanie VCC (powinno być 3,3 V).
RuntimeError: Zły ACK: [lista bajtów]	MOSI i MISO prawdopodobnie zamienione miejscami. Sprawdź tabelę połączeń.
Układ: PN500 lub inne dziwne wartości	Problem z linią danych – sprawdź, czy MISO jest poprawnie podłączone.
Thonny nie widzi Pico	Sprawdź kabel USB (niektóre kable ładują tylko, nie przesyłają danych). Spróbuj innego.
Brak wyjścia po uruchomieniu main.py	Upewnij się, że wgrałeś oba pliki NA PICO, nie lokalnie na dysk komputera.

Bonus: prosty wykrywacz kart RFID

Mamy już działającą komunikację z PN532. Możemy szybko sprawdzić, czy układ rzeczywiście widzi karty – bez pisania pełnej obsługi protokołu ISO 14443 (to dopiero następny odcinek). Poniższy listing korzysta bezpośrednio z niskopoziomowych metod klasy PN532, żeby wysłać komendę InListPassiveTarget i sprawdzić, czy w polu czytnika jest karta.

Wbudowana dioda LED Pico (GPIO 25) miga przy każdym wykryciu:

Listing 3: wykrywacz_pola.py

```
# wykrywacz_pola.py
# Wykrywanie kart RFID – dioda LED miga przy każdym znalezieniu

from machine import SPI, Pin
from pn532 import PN532
import time

spi = SPI(0, baudrate=1_000_000, polarity=0, phase=0,
          sck=Pin(18), mosi=Pin(19), miso=Pin(16))
czytnik = PN532(spi=spi, cs_pin=17)

led = Pin(25, Pin.OUT) # wbudowana dioda LED Pico

ic, ver, rev, _ = czytnik.get_firmware_version()
print(f'PN532 gotowy (firmware {ver}.{rev}). Czekam na kartę...')

CMD_RF_CONFIGURATION = 0x32
CMD_IN_LIST_PASSIVE_TARGET = 0x4A

# Włącz pole RF
czytnik._send_command(CMD_RF_CONFIGURATION, bytes([0x01, 0x01]))

while True:
    try:
        czytnik._send_command(CMD_IN_LIST_PASSIVE_TARGET,
                              bytes([0x01, 0x00])) # 1 karta, ISO14443A
        data = czytnik._read_response(CMD_IN_LIST_PASSIVE_TARGET, 20)
        if data[0] > 0: # data[0] = liczba znalezionych kart
            led(1)
            uid_len = data[5]
            uid = data[6:6 + uid_len]
            print('Karta:', ' '.join(f'{b:02X}' for b in uid))
            time.sleep_ms(500)
            led(0)
    except Exception:
        pass
    time.sleep_ms(200)
```


Uwaga: po co _send_command z podkreślnikiem?

W listingu 3 używamy prywatnych metod klasy PN532 bezpośrednio. Robimy to tutaj tylko po to, żeby pokazać, że PN532 już teraz potrafi wykrywać karty.

W rozdziale 2 napiszemy klasę ISO14443, która zamknie te wywołania za czystym, publicznym interfejsem. Od tamtego momentu nie będziemy już dotykać metod z podkreślnikiem z zewnątrz.

Wgraj plik na Pico, uruchom i przykładaj różne karty do modułu. W konsoli powinieneś widzieć kolejne linie z heksadecymalnymi UID kart:

```
PN532 gotowy (firmware 1.6). Czekam na kartę...
Karta: 04 A3 1B 7C
Karta: 04 A3 1B 7C
Karta: 92 F1 04 B3
```

Każda karta ma swój unikalny UID. Przyłóż tę samą kartę kilka razy – UID będzie taki sam. Przyłóż inną kartę – UID się zmieni. To dobry test poprawności sprzętu przed przejściem do kolejnego rozdziału.

Podsumowanie

W tym rozdziale zrobiliśmy całkiem sporo. Dowiedzieliśmy się, jak różnią się trzy

rodziny RFID i dlaczego w tym kursie skupiamy się na 13,56 MHz. Zrozumieliśmy, jak karta może działać bez baterii i jak przebiega transmisja danych bez własnego nadajnika po stronie karty. Podłączyliśmy moduł PN532 do Raspberry Pi Pico przez SPI, rozkodowaliśmy format ramki protokołu PN532 i napisaliśmy pierwszą klasę – warstwę 1 naszej docelowej architektury.

Klasa PN532 w pliku pn532.py będzie towarzyszyła nam przez cały kurs. W każdym kolejnym rozdziale będziemy nad nią budować nową warstwę, ale samej klasy nie będziemy już ruszać – i właśnie o to chodzi w warstwowym podejściu do kodu.

W rozdziale 2 napiszemy klasę ISO14443, która zamieni surowe bajty z PN532 w sensowne obiekty: karty z UID, SAK i ATQA. A przy okazji dowiemy się, jak dokładnie działa pętla antykolizyjna opisana w normie ISO/IEC 14443-3.

Pracownia Konstrukcyjna AVT

Materiały do dalszej lektury

NXP PN532 User Manual (UM10232) – dostępny bezpłatnie na nxp.com. Rozdział 6 opisuje protokół, rozdział 7 – pełną listę komend. Elechouse PN532 NFC RFID Module V3 User Guide – schemat, opis przełączników, przykłady podłączeń. Do pobrania ze strony elechouse.com lub z GitHub Elechouse. ISO 14443 Protocol Guide: Types, Tech & Smart Card Uses – przystępny artykuł wprowadzający w warstwy normy ISO 14443 (łatwo znaleźć przez Google).

REKLAMA

Altium® Agile

Gdy zespoły mówią
jednym językiem,
projekty nabierają tempa.

Usprawnij współpracę między
elektroniką, mechaniką i produkcją
w jednym środowisku projektowym.



**LEPSZA KOMUNIKACJA
MIĘDZY DZIAŁAMI**

Płynna współpraca
w czasie rzeczywistym.



**JEDNO ŹRÓDŁO
PRAWDY**

Wspólne dane.
Mniej błędów.



**PEŁNA KONTROLA
I BEZPIECZEŃSTWO**

Bezpieczne
zarządzanie dostępem
i procesami.



**UPROSZCZONE
PROCESY**

Szybsze decyzje.
Płynna realizacja.



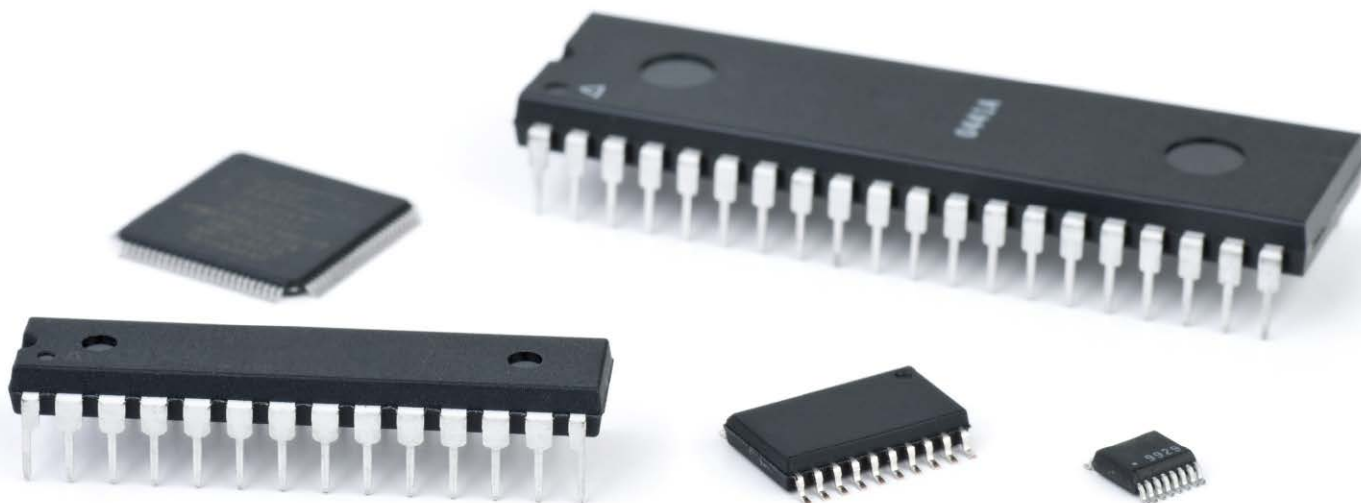
**PEŁNA WIDOCZNOŚĆ
PROJEKTU**

Od koncepcji
po produkcję.

**COMPUTER
CONTROLS**

Altium Agile Teams to zintegrowana
platforma dla zaawansowanej współpracy
multidyscyplinarnej.

www.ccontrols.pl
Oficjalny dystrybutor Altium



„Jaki mikrokontroler wybrać do ...”, czyli przypowieść o Wojtku

W ostatnich miesiącach trafiłem na serię filmów na YouTube, które łączyły dwie rzeczy: osobę autora oraz opinie na temat najlepszych i najgorszych rodzin mikrokontrolerów, z którymi to opiniami się nie zgadzam. I nie jestem w tym osamotniony. Nie podam nazwiska rzeczowego twórcy, ani nie podlinkuję jego kanału, gdyż moim zdaniem nie zasługuje on na jakąkolwiek formę promocji, więc dajmy mu na imię Wojtek.

Wojtek każdy film zaczyna od stwierdzenia, iż jest ekspertem, bo pracował dla dużego amerykańskiego producenta układów scalonych, ale od niego odszedł i wypuścił własny produkt. Takie uporczywe podkreślanie własnych osiągnięć w każdym filmie wzbudza moją podejrzliwość.

Fachowcy nie muszą się chwalić swoją fachowością – to zwyczajnie widać w tym, co robią lub prezentują. Dave Jones z kanału EEVBlog nie zaczyna każdego swojego filmu od listy osiągnięć. Moim prywatnym, osobistym zdaniem kariera Wojtkowi nie do końca wyszła, dlatego teraz nagrywa filmy na YouTube.

W filmach tych wyraża opinie o wątpliwej wartości merytorycznej, popełniając przy tym szkolne błędy. Od kiedy to PIC, AVR czy STM32 to nazwy architektur? (Autor zmienił tytuł filmu na bardziej poprawny po zwróceniu uwagi przez widzów.) Film z tym błędem skłonił mnie do napisania tego felietonu.

Gwoli wyjaśnienia – te ośmiobitowe mikrokontrolery AVR i PIC to układy RISC oparte na zmodyfikowanej architekturze harwardzkiej.

Rodzina 16-bitowych PIC-ów też używa architektury harwardzkiej RISC, zoptymalizowanej dla języka C, a rodzina PIC32 opiera się na rdzeniach MIPS M4K/M-Class oraz ARM Cortex-M.

Z kolei rodzina STM32 od STMicroelectronics to rdzenie ARM Cortex-M. Wojtek w ogóle odradza używanie mniej niż 32 bitów. O architekturze '50 czy 68k nawet nie wspomina.

Dla niego najlepsze mikrokontrolery to te, które mają Wi-Fi i Bluetooth Low Energy, czyli właśnie wybrane modele STM32, ESP32 czy układy od firmy Nordic.

Dla Wojtki może to być szok, ale nie każde urządzenie potrzebuje procesora ARM ani łączności radiowej z całym światem. Ba, architektura '51, która ma więcej lat

niż ja w 2024 roku, miała wartość rynkową 8,5 miliarda dolarów i będzie rosła dalej.

W omawianym przeze mnie w poprzednim wydaniu EP module MaixCAM jednym z rdzeni był właśnie 8-bitowy 8051, choć nie był on bezpośrednio dostępny dla programisty.

Dla Wojtka (i dla Czytelnika) szokującym może być fakt obecności na rynku mikrokontrolerów 4-bitowych, których wartość rynkowa w 2024 roku wyniosła 2,1 mld dolarów.

Układy te trafiają na przykład do szczoteczek elektrycznych Braun Oral-B.

W sumie nie potrzeba zbyt wiele mocy obliczeniowej, by odliczać czas i sterować prostym silnikiem elektrycznym i diodą LED oraz kontrolować napięcie akumulatora i stan przycisku zasilania.

W innych filmach Wojtek popełnia podobne błędy. Widać niezdrową fascynację ARM i RISC-V z odrzucaniem wszystkiego, co było wcześniej lub co jest inne.

Dziwi na przykład krytyka mniej znanych chińskich producentów mikrokontrolerów kosztujących centy, gdy jednocześnie Wojtek zachwyci się rodziną ESP32, również chińskiego producenta (ale już nie zaleca ESP8266).

Pomija fakt, iż Espressif Systems sprawiło szpetną niespodziankę użytkownikom przy przejściu ESP-IDF z wersji 4.x na 5.x, co spowodowało zmiany w API i wycofanie niektórych starych funkcji, z których korzystali użytkownicy.

O ile faktycznie chińskie mikrokontrolery mniej znanych producentów mają duże problemy z dobrą dokumentacją czy z IDE, to cena liczona w centach za układ do prostych zastosowań jest nie do pobicia.

W innym filmie Wojtek wymienia mikrokontrolery do zastosowań bateryjnych. Z jakiegoś powodu nie ma nic o PICach, za to na pierwszym miejscu jest rodzina Ambiq Apollo 510. Co to za rodzina i producent? Nie mam zielonego pojęcia, bo pierwszy raz o nich słyszę.

W polskich sklepach nie do kupienia, u zagranicznego dystrybutora przykładowy układ kosztuje 40 złotych od sztuki. Pobór prądu? Trudno powiedzieć, bo na początku noty jest podana wydajność na milidżul. W tabeli poboru prądu mamy wyniki w mikrowatach. Trudno powiedzieć, czy się poborem chwalam, czy się go wstydzę.

Swoją drogą rodzina Apollo 510 to dość potężna rodzina układów z GPU i wsparciem dla sieci neuronowych. Wojtek zaleca go do zastosowań bateryjnych, tylko nie podaje takich faktów, jak pobór prądu w głębokim uśpieniu – 13,8 μA /1,8 V.

Układy PIC32CM Lx osiągają 1,7 μA w trybie standby z podtrzymaniem całej pamięci RAM i poniżej 100 nA w trybie wyłączonym.

Wojtek w filmie o najlepszych układach do zastosowań bateryjnych wymienia też układy TI z rodziny MSP430, wymienia je też jako jedne z najgorszych układów do nowych

projektów i jako jedną z najgorszych „architektur” (choć zmienił tytuł filmu na „rodziny” po moim upomnieniu).

Jeden Wojtek, kilka filmów, a tyle już napisałem. Zmieńmy więc nieco temat i porozmawiajmy o sensie rekomendowania konkretnych rodzin, architektur czy marek oraz po czym można poznać złego inżyniera.

Wybory, wybory...

Pytanie o wybór mikrokontrolera czy rodziny mikrokontrolerów do projektu jest o tyle trudne, iż do wyboru mamy tysiące układów od dziesiątek producentów.

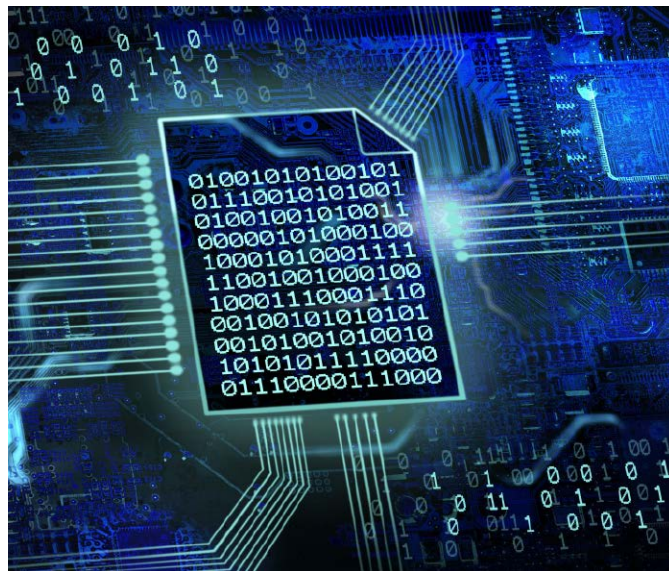
Większość układów w produkcji to wariacje na temat kilku bazowych typów.

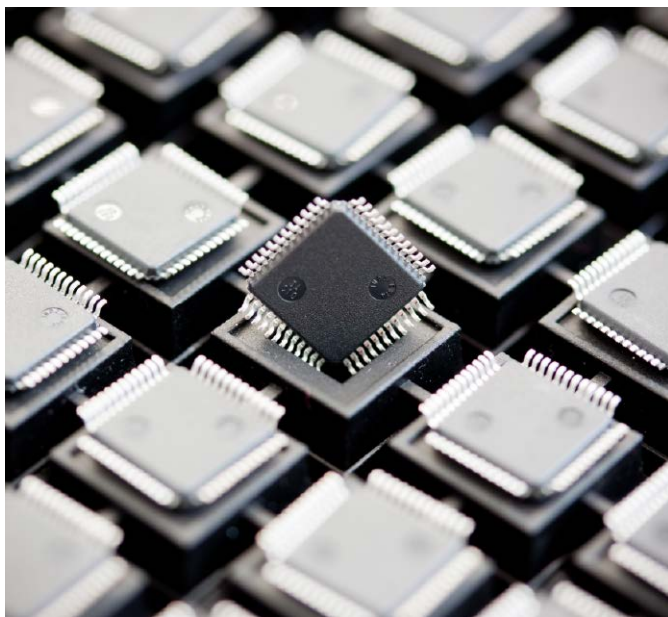
Generalnie największą różnicą między nimi jest liczba bitów w ALU, co przede wszystkim wpływa na wydajność obliczeniową, pobór prądu i złożoność programowania, jeśli wzorem naszych przodków będziemy pisać programy w assemblerze. Ale ponieważ wszyscy piszą już w C/C++, to architektura nie ma znaczenia dla programisty. Co się zatem liczy?

Peryferia, pobór energii i cena. Oraz dostępność w hurcie. Kilku komentujących film Wojtka o najgorszych architekturach wspomniało, iż stworzyli produkty na układach PIC, które po dekadzie lub dwóch wciąż są produkowane i sprzedawane w setkach tysięcy egzemplarzy. Prawda jest taka, że w większości urządzeń codziennego użytku osiem lub mniej bitów w zupełności wystarczy.

Blender czy mikrofalówka nie potrzebują niczego lepszego niż PIC12F/16F czy STM8, by obsłużyć załączanie i wyłączanie zasilania, regulację mocy czy odmierzenie czasu.

Mikrokontroler STM8 znalazłem w najlepszym podgrzewaczu do mleka i jedzenia dla dzieci, jaki był dostępny na rynku (nie jest produkowany od lat i sam musiał kupić model używany sześć lat temu). Urządzenie nie miało nawet czujnika temperatury (tylko bezpiecznik termiczny





dla grzałki) i dobierało czas podgrzewania na bazie wskazanego rodzaju produktu (mleko czy jedzenie), objętości (co bodaj 10 ml) oraz tego, czy pojemnik jest w temperaturze pokojowej, wyjęty z lodówki czy z zamrażarki. Jedną wadą był zasilacz beztransfornatorowy, który się z czasem psuł, oraz nierozbieralna konstrukcja.

Dobłą zasadą wyboru mikrokontrolera do projektu jest sprawdzanie, co spełni minimalne wymagania pod kątem peryferiów, ceny i dostępności. Jeśli to nie jest projekt IoT, to nie trzeba mikrokontrolera z wbudowanym radiem.

Nawet jak to jest projekt IoT, to może się okazać, że taniej i prościej będzie użyć gotowego modułu radiowego: Wi-Fi, BLE czy LoRa.

Odpada wtedy konieczność certyfikacji radia pod względem EMI/EMC. Często moduły takie mają lepsze parametry pod względem poboru prądu, zwłaszcza w trybie uśpienia.

W wielu przypadkach nie potrzeba też dużej wydajności mikrokontrolera. Apollo Guidance Computer, który pozwolił Amerykanom wylądować na Księżycu, posiadał zegar 1,024 MHz, 16-bitową architekturę, ale z powodu unikalnej konstrukcji jego wydajność wynosiła tylko maksymalnie 0,085 MIPS. A mimo to wystarczył do zadań, do których został stworzony.

Inaczej pisząc, nie potrzeba układu taktowanego zegarem 240 MHz z dwoma wysoce wydajnymi rdzeniami ARM, by obsłużyć prosty termostat albo kontrolować garść programowalnych LED-ów.

Przed erą mikrokontrolerów termostat robiło się z tranzystora, rezystora, termistora i potencjometru. Jeszcze wcześniej wystarczyły dwie blaszki z różnych metali i wkret do regulacji odległości styku od końca tych blaszek. Swoją drogą takie przełączniki termiczne są wciąż używane, głównie jako zabezpieczenia.

W większości sytuacji moc obliczeniowa mikrokontrolera nie ma tak naprawdę znaczenia, gdyż każdy mikrokontroler

jest w stanie wykonać każde zadanie, pod warunkiem, iż ma wystarczająco dużo miejsca na program. Nawet nie trzeba zbyt wiele pamięci RAM, bo na rynku dostępne są układy SRAM komunikujące się przez interfejs szeregowy.

Czy ma sens ładowanie modelu AI do ośmiobitowego mikrokontrolera? Oczywiście, że nie, ale moim zdaniem generalnie nie ma zbyt wielu sytuacji, w których trzeba ładować model AI do systemu embedded.

Kwestia peryferiów jest dość ciekawa. Microchip oferuje setki różnych mikrokontrolerów ośmiobitowych, które różnią się głównie pojemnością pamięci programu, RAM i Flash/EEPROM oraz właśnie kombinacją peryferiów. Najbardziej „wypasione” mają wszystkiego po trochu, ale generalnie można dobrać coś pod konkretną sytuację.

Ponad dekadę temu trafiłem na przykład na układ przeznaczony do budowy przetwornic impulsowych – i to nie jakichś prostych bucków czy boostów, ale wszystkich topologii, którym wystarczą cztery sygnały kontrolne.

Wielu „starych wygów” mówiło mi wtedy, że nie da się zrobić przetwornicy na mikrokontrolerze, zwłaszcza takiej pracującej z dużymi mocami, bo będzie się ona wieszać od EMI i trzeba czegoś analogowego w stylu TL494.

Najwidoczniej ktoś tego nie powiedział inżynierom z Microchip, bo ci wzięli i zaprojektowali PIC16F785 właśnie do budowy zasilaczy i nawet przetestowali kilka przykładowych konfiguracji.

Mój pomysł na ten układ zakładał budowę zasilacza, który składa się z „klocków” zależnie od potrzeb i konfiguruje się przez UART.

Projekt okazał się trochę zbyt ambitny jak na moje ówczesne umiejętności, ale nauczyłem się sporo o przetwornicach i o konieczności prowadzenia dobrej dokumentacji.

Innym moim ambitnym projektem był emulator PDP-8 na PIC24. Dlaczego? Bo jedyny projekt emulatora jakiegokolwiek komputera DEC PDP potrzebuje Raspberry Pi – komputera 20...30 tysięcy razy wydajniejszego od najszybszego wariantu PDP-11/70, który emuluje. Z wybranym układem PIC24 osiągnąłbym prędkość zbliżoną do PDP-8/S, co by mi w zupełności wystarczało. Teraz prawdopodobnie wybrałbym inny model układu, ale też raczej z rodziny PIC24/dsPIC, zwracając szczególną uwagę na grubość erraty. Ale moim skrytym marzeniem jest emulator komputera Odra 1305 na FPGA.

Wracając do mikrokontrolerów – to istnieje wiele robiących wrażenie projektów, które są realizowane na Arduino, czyli ATmega328 z IDE, które, prowadząc początkującego programistę „za rączkę”, ma raczej mało wydajne funkcje dostępu do I/O i peryferiów. Programista nieco bardziej zaawansowany ma możliwość sięgnąć do „gołego metalu” i operować bezpośrednio na rejestrach.

Ostatnia kwestia to cena i dostępność. Weźmy na przykład płytkę rozwojową Black Pill z układem STM32F411CEU6.

Na AliExpress takie płytki kosztują 7...15 złotych od sztuki. Sam mam ze dwie i programowałem je w VS Code z pomocą PlatformIO. Ale jeśli chciałbym zrealizować komercyjny projekt na tym układzie ARM od STMicroelectronics, to u polskiego dystrybutora sam układ kosztuje 25,50 zł netto. Jak to możliwe, skoro płytka z Chin z tym układem kosztowała sporo mniej? Odpowiedź jest prosta: polski dystrybutor komponentów elektronicznych ma oryginalne układy z fabryki STMicroelectronics, a w Chinach mają klony, które mogą nie spełniać parametrów podanych w nocie katalogowej oryginału.

To mi przypomina ciekawy przypadek układów nRF24L01+ firmy Nordic – oryginalne układy miały błąd w krzemie, przez co nie wszystko działało poprawnie. Chińskie klony tych układów powstały częściowo na podstawie noty katalogowej i tego błędu nie mają.

Wybierając mikrokontroler do masowej produkcji (oraz inne komponenty) trzeba patrzeć na cenę hurtową. Przy okazji może się okazać, iż taniej wyjdzie wybrać nieco droższy mikrokontroler, który ma na przykład wbudowany wzmacniacz operacyjny rail-to-rail, niż kupować ten komponent osobno. Przy okazji można też oszczędzić miejsce na PCB. Kwestia dostępności wygląda nieco ciekawiej.

Układy od Microchip czy STMicroelectronics mogę kupić zarówno od polskiego dystrybutora, jak i od kilku dystrybutorów zagranicznych. Ba, mogę zamówić te układy bezpośrednio od producenta, z opcją wgrania firmware w fabryce, a za dodatkową opłatą mogę dostać własne logo i nazwę na układzie, by utrudnić inżynierię wsteczną mojego nowego produktu.

W jednym Wojtek miał rację: kupując najtańszy chiński mikrokontroler jesteśmy zależni od chińskiego producenta i chińskiego dystrybutora i zazwyczaj nie ma alternatyw.

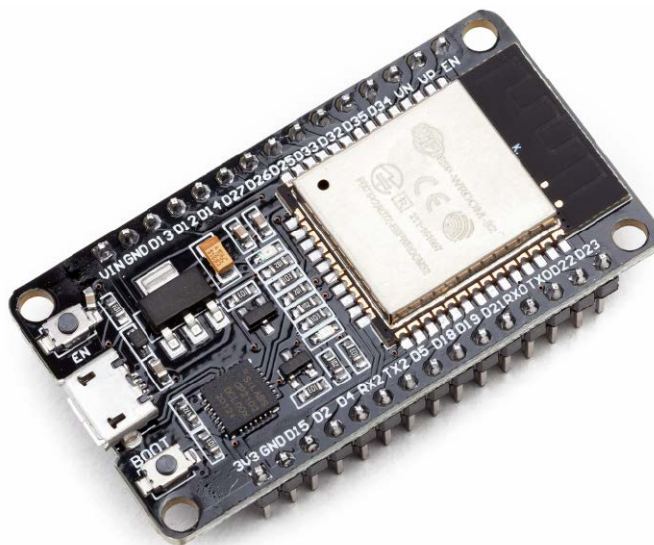
Biorąc pod uwagę obecną, niestabilną sytuację geopolityczną oraz wchodzące niedługo cła na chińskie produkty, stosunek ryzyka do oszczędności zaczyna wyglądać mniej korzystnie dla Chin. Zachodni producenci zazwyczaj są w stanie zagwarantować długoterminową dostępność komponentów.

Złej baletnicy...

Niektórzy ludzie, głównie hobbysci, skupiają się nadmiernie na mało istotnymi kwestiami.

Jedni nie lubią MPLAB X IDE, inni mają problem z STM32CubeIDE. Jeszcze inni narzekają na IDE od Espressif.

Przez ostatnie 30 lat z kawałkiem używałem wielu środowisk programistycznych, wielu narzędzi do projektowania CAD/CAM, wielu narzędzi do projektowania czy symulowania obwodów, a także narzędzi do modelowania 3D, montażu wideo, VFX, edycji obrazów czy produkcji muzycznej. Używałem DOS-a, Windowsa 3.11, 98SE,



ME, XP, Vista, 7 czy w końcu 10/11. Testowałem Linuksy od Damn Small Linux przez Knoppix, kilka generacji i odmian Ubuntu, Red Hat czy openSUSE.

Znam języki C/C++, Pascal (i Delphi), BASIC, wliczając w to warianty C64, +4 i Atari 800XE, LOGO, Python, wariant języka Java używany w grze Colobot, VHDL oraz Lua.

I czego mnie to nauczyło? Zasadniczo jednej rzeczy: nie jest ważne ani IDE, ani język programowania. Ważna jest wiedza na temat algorytmów i tworzenia logiki programu oraz umiejętność czytania dokumentacji. Niektóre IDE są lepsze niż inne, podobnie jak języki programowania. Dla przykładu w pogardzie i nienawiści do Pythona zostałem wychowany. I każdego języka, który używa niewidocznych znaków jako części składni.

Jeśli IDE pozwala napisać kod źródłowy, skompilować go i wgrać do docelowego układu, to spełnia minimalne wymagania. Dobrze, jak pozwala też debugować kod w układzie i podglądać rejestry. Bardzo dobrze, jak ma przyzwoite biblioteki standardowe i rozbudowaną dokumentację z przykładami.

Konfiguratorzy czy generatory kodu to dodatek, z którego nie miałem powodu korzystać (poza ustawianiem rejestrów konfiguracyjnych), ale jeśli tylko generowany kod jest łatwy w użyciu, to nie widzę powodu, by nie uznać ich za zaletę. Wygląd środowiska programistycznego w ogóle nie ma znaczenia, jeśli tylko najistotniejsze funkcje da się łatwo zlokalizować.

Serio, pisałem funkcjonalny kod w środowisku wziętym żywcem z DOS-a. Jeśli mogę z IDE korzystać pomimo wady wzroku dzięki funkcjom ułatwień dostępu, to nie obchodzi mnie, że wygląda, jakby powstało w 1994 roku na Windows 3.11.

Jeszcze 10...20 lat temu było sporo dyskusji pod zbiorczym hasłem „Jaki język programowania mikrokontrolerów jest najlepszy?”. U nas dzięki kursom w EdW popularny był Bascom i układy AVR. W USA był BASIC Stamp na PIC-u. W ogóle układy PIC były tak popularne w USA, jak u nas



AVR. Najczęściej programowane w assemblerze, bo był bezpłatny, a i nie było zbyt wielu alternatyw.

U nas assembler też był popularny, choć raczej dla układów '51, bo te były (i chyba wciąż są) obecne na uczelniach.

Potem pojawiło się Arduino, czyli C/C++ z ułatwieniami dostępu. Ja wtedy odkrywałem uroki mikroPascala i PICBasica, zanim przerzuciłem się na C. To jest, moim zdaniem, najlepszy wybór dla mikrokontrolerów niezależnie od rodziny, liczby bitów i prędkości.

Styl programowania funkcjonalny: program rozbity jest na funkcje, z których każda wykonuje specyficzne zadanie, funkcje mogą wywoływać inne funkcje i przekazywać im parametry. Ten styl programowania sprawia najmniej problemów zarówno początkującym, jak i bardziej zaawansowanym.

Nie polecam programowania obiektowego – niepotrzebnie komplikuje proste rzeczy. W końcu nie tworzymy systemu operacyjnego albo zaawansowanej aplikacji takiej jak pakiet biurowy, po co więc komplikować sobie życie?

Rozbijając program na proste funkcje, możemy skupić się na poszczególnych elementach i łatwiej go zmodyfikować czy rozbudować w przyszłości.

Warto też rozważyć stosowanie prostych maszyn stanów skończonych z wykorzystaniem zmiennej „volatile” do przechowywania bieżącego stanu programu oraz instrukcji warunkowej „switch-case” w pętli głównej do wywoływania kolejnych funkcji według potrzeby. Program staje się dzięki temu bardziej modułowy, bo nowe funkcjonalności dodaje się, dopisując odpowiednie funkcje oraz kolejny blok „case” w pętli głównej. Kompilator i tak to zoptymalizuje. Warto też komentować i dokumentować

kod oraz pamiętać o sensownych nazwach dla funkcji, stałych i zmiennych.

Czytelniku, jeśli ktoś Ci mówi, że kod źródłowy powinien być swoją własną dokumentacją, to zalecam trzymać się z daleka od takiej osoby. Po co komuś utrudniać życie, by musiał zgadywać, co i jak dana funkcja robi, skoro można dopisać parę zdań tu i ówdzie? Dokumentowanie ma też tę zaletę – łatwiej pilnować, by kod robił to, co ma, i aby nie dodawać nadmiaru zbędnych funkcji (feature creep).

Moje pierwsze programy nie miały dobrej dokumentacji i gdy do nich wracałem po pół roku, to nie rozumiałem wcale, dlaczego coś było napisane tak, jak było napisane i co oznaczają różne magiczne liczby, którymi kod był upstrzony, jak pomnik na starówce ptasimi odchodami.

Niektórzy twierdzą, iż Python wymusza czytelność kodu przez używanie klawisza [Tab] jako części składni.

Przez ostatnie 15 lat nie widziałem środowiska programistycznego, które nie potrafiłoby automatycznie dodawać wcięć. Ba, większość potrafi „zwijać” funkcje, by zajmowały mniej miejsca, gdy nie trzeba w nich grzebać.

Ja wiem, że w czasach, gdy Python powstawał, taka magia była nie do pomyślenia, no ale jesteśmy za połową trzeciej dekady XXI wieku i czas dorosnąć.

Jak ktoś chce koniecznie bawić się w języki interpretowane, to zdecydowanie Lua, FORTH, a nawet leciwy BASIC są lepszym wyborem niż wyroby pythonopodobne.

Jasne, można i w Pythonie programować, jak się nie zna niczego lepszego, ale płaci się wydajnością i wymogiem używania układu 32-bitowego.

Jest jednak wielka zaleta układów w rodzaju ESP32 czy płytek modułów opartych na Arduino, STM32 lub RP2040 – nie potrzeba programatora.

Dla hobbysty to często spora oszczędność – programator PICKit Basic kosztuje około 150 złotych, a PICKit 5 to już koszt ponad 350 złotych.

Jeśli projekt wymaga Wi-Fi, to moduły ESP8266 i ESP32 wydają się świetnym wyborem, gdyż odpada przy okazji problem certyfikacji urządzenia radiowego, skoro sam moduł ma wszystkie potrzebne certyfikaty zgodności.

To właśnie sprawiło, że na rynku od dłuższego czasu jest sporo gniazdek, ściemniaczy światła i innych urządzeń IoT, w których wewnątrz pracuje moduł oparty na ESP8266.

Gdybym potrzebował aż takiej mocy obliczeniowej w swoich projektach, to używałbym w nich modułów opartych na płytkach BlackPill z układami STM32 ze względu na łatwość ich użycia.

W produkcji seryjnej zawsze można wykorzystać „goły” układ zaprogramowany w trakcie montażu. Nie zawsze jednak można użyć gotowego modułu czy nawet układu z tego modułu; kilka lat temu widziałem otwarty projekt kontrolera silnika (ECU) opartego na Arduino. Jako

projekt edukacyjny jest to jeszcze w miarę dobre rozwiązanie. Jako alternatywa dla fabrycznego modułu ECU jest to skrajny idiotyzm.

Arduino bowiem, i to oryginalne, nie spełnia wymogów branży motoryzacyjnej, o czym sami twórcy Arduino informują. A większość ludzi używa chińskich podróbek, które mogą być gorszej jakości. To kolejna rzecz, która odróżnia dobrego elektronika czy programistę od złego: wie, jakich narzędzi, układów czy modułów użyć w danym projekcie. Hobbysta może robić, co i jak chce, ale jeśli coś ma iść do produkcji, to powinno być zrobione poprawnie i na właściwych komponentach.

Wiele projektów na Kickstarterze kończyło marnie, bo nie „dowodziły” tego, co obiecano, przy czym zwykle z powodu błędów projektowych już na wstępnym etapie prototypowania.

Hobbysta często brnie w stosowanie niewłaściwego rozwiązania i próbuje przymusić je do działania, zamiast wybrać coś lepszego, bo zainwestował w nie czas i pieniądze.

Ta pułapka utopionych kosztów (ang. sunk cost fallacy) jest przyczyną porażki niejednego projektu, ale też zdarza się i w życiu codziennym, od prób ciągłego naprawiania gruchota, który powinien już dawno temu skończyć na złomie, po próby ratowania związku z osobą, która już dawno temu powinna być rzucona.

Projektant musi wiedzieć, że jeśli dane rozwiązanie nie będzie działać mimo kilku iteracji projektu, to trzeba z niego zrezygnować i wybrać inne. Sam się tak wkopałem w jednym, małym projekcie, wybierając „na upartego” rozwiązanie, które było nieoptymalne.

Użyłem bowiem taniego wzmacniacza operacyjnego i zbudowałem dwie pompy ładunku (jedną do podniesienia napięcia zasilania, drugą do odwrócenia polaryzacji dla gałęzi ujemnej), taktowane sygnałem CLKOUT mikrokontrolera, zamiast zwyczajnie dopłacić te dwa złote i wybrać inny wzmacniacz operacyjny.

Moje rozwiązanie jak najbardziej działało, ale z punktu masowej produkcji było absolutnie bezsensowne.

Zakończenie

Na pytanie „Jaki mikrokontroler wybrać do...?” odpowiedź jest prosta: najtańszy i najprostszy, który wykona powierzone mu zadanie, niezależnie od architektury, IDE i języka programowania (choć C/C++ jest preferowanym wyborem).

Podobnie z wyborem pozostałych komponentów. Warto też wiedzieć, kiedy konieczna może być zmiana wybranego rozwiązania, zanim utopi się w nim zbyt wiele pieniędzy.

Warto też nie ufać Wojtkom z YouTube, niezależnie od tego, gdzie pracowali i co robili.

Paweł Kowalczyk



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)

