

ELEKTRONIKA PRAKTYCZNA

● MIĘDZYNARODOWY MAGAZYN ELEKTRONIKÓW KONSTRUKTORÓW ● WRZESIEŃ ● 9/2026 ●



FOTOWOLTAIKA BALKONOWA

Z praktyki ekspertów

- Profesor Christophe Basso. Stabilizacja pętli trójfazowego prostownika Vienna (2)
- Ochrona przepięciowa układów (1)

W pracowni

- Czujniki wilgotności i cieczy dla domu i ogrodu
- Druk 3D w służbie elektroniki (11)

Repetitorium

- Obudowy w elektronice profesjonalnej
- Zaawansowane funkcje oscyloskopów cyfrowych dla niezaawansowanych (1)
- PCB w praktyce, część 1. Projektowanie PCB

Kity AVT

- Włącznik pomocniczy (AVT6102)

AI w elektronice

- AI w pracowni elektronika konstruktora (4)
- Kod z autopilota. Programowanie mikrokontrolerów wspomagane przez AI

Kursy

- Kurs RFID z PN532.
 - Rozdział 5. MIFARE Ultralight i EV1: pamięć bez Crypto1
 - Rozdział 6. NTAG216: mirror UID i licznika – dynamiczne tagi bez zapisu
- Laboratorium pomiarowe dla konstruktora. Oscyloskop 12-bit w praktyce konstruktora: ile bitów naprawdę trafia na ekran
 - Rozdział 1. 12 bitów to nie zawsze 12 bitów – co rozdzielczość i tryb High-Res dają konstruktorowi, a czego nie załatwią
 - Rozdział 2. Wyzwalanie i pomiary automatyczne: ukryta inteligencja oscyloskopu

E-prenumerata Elektroniki Praktycznej

Tvoja 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

Fotowoltaika balkonowa

W Polsce pracuje dziś około 1,6 miliona domowych instalacji fotowoltaicznych o łącznej mocy przekraczającej 14 GW. To jedna z największych flot prosumenckich w Europie i dorobek niespełna dekady. Zaczynało się od pojedynczych instalacji entuzjastów, a skończyło na tym, że panele na dachu domu jednorodzinnego są dziś widokiem zupełnie zwyczajnym. Jest jednak grupa obywateli, której ten sukces właściwie nie dotyczy. Prosumentem zostaje się u nas razem z własnym dachem. Mieszkaniec bloku dachu nie ma i mieć nie będzie. Ma natomiast balkon, a na balustradzie mieszczą się dwa moduły. Z mikrofalownikiem ograniczonym do 800 W daje to instalację, która kosztuje mniej więcej tyle co przyzwoity rower i nie wymaga ani ekipy montażowej, ani rusztowania.

W Europie ten pomysł już się przyjął. W Niemczech zarejestrowanych jest około 1,4 miliona takich zestawów o łącznej mocy modułów rzędu 1,5 GWp, a w samym ubiegłym roku przybyło ich tam blisko pół miliona. Kupuje się je w marketach obok wiertarek. Podobny ruch widać w Austrii, Holandii, Hiszpanii i we Włoszech. Przełomem nie była żadna nowa technologia – moduł fotowoltaiczny i mikrofalownik wyglądają tak samo jak pięć lat temu. Przełomem było uproszczenie przepisów: podniesienie limitu mocy i zdjęcie z użytkownika większości formalności.

W Polsce zainteresowanie jest duże. Widać je po ruchu w sklepach, po dyskusjach na forach, po pytaniach, które trafiają do naszej redakcji. Bariera nie leży już w technice ani w cenie. Leży w zgodach wspólnot i spółdzielni oraz w tym, że zestaw o mocy 800 W traktujemy formalnie dokładnie tak samo jak instalację dachową o mocy 10 kW.

Ciekawe jest przy tym, że pytania czytelników rzadko zaczynają się od okresu zwrotu. Motywacje bywają bardzo różne.

Dla części osób najważniejsza jest niezależność. Nie chodzi o zerwanie z siecią, bo przy 800 W nie ma o tym mowy, tylko o posiadanie własnego źródła, nad którym ma się kontrolę. Innym chodzi o bezpieczeństwo. Wielogodzinna awaria zasilania na Półwyspie Iberyjskim pokazała, że w Europie nadal potrafi zgasnąć światło w całym kraju, a magazyn

energii przy balkonowej instalacji utrzyma przez ten czas router, komputer, światło LED, lodówkę i ładowarkę do telefonu. To całkiem komfortowe warunki dla przetrwania sytuacji awaryjnej.

Są też motywacje, których nie da się przeliczyć na złotówki. Satysfakcja z uruchomienia czegoś, co działa i co się samemu zrozumiało – dla czytelników EP argument chyba najbardziej zrozumiały. Ciekawość, ile właściwie pobiera lodówka i co się dzieje w mieszkaniu o trzeciej w nocy, bo instalacja z pomiarem prądu pokazuje to natychmiast. Chęć zrobienia czegoś sensownego dla środowiska, w skali dostępnej pojedynczej osobie. Wreszcie zwykła przyjemność majsterkowania i to w dziedzinie, w której efekt widać na rachunku za prąd.

Zmieniło się przy tym coś istotnego w samej technice. Zestaw balkonowy bez magazynu energii oddaje w środku dnia większość produkcji do sieci, bo w mieszkaniu nikogo wtedy nie ma. Małe magazyny energii, które pojawiły się przy tych instalacjach, długo działały źle, bo nie wiedziały, ile mieszkanie zużywa energii w danej chwili. To już przeszłość. Artykuł okładkowy w tym numerze pokazuje, jak działają mikroinstalacje fotowoltaiczne z magazynami i jak wygląda ich rachunek ekonomiczny w polskich warunkach.

Wniosek z tego wszystkiego jest jeden. Fotowoltaika balkonowa nie czeka na wynalazek, tylko na przepisy i na zmianę przyzwyczajzeń zarządców budynków. Technika jest gotowa, ceny są względnie przystępne, a chętnych nie brakuje. Szkoda byłoby zostawić trzecią część mieszkańców naszego kraju poza rynkiem, na którym reszta obywateli radzi sobie tak dobrze.

Wiesław Marciniak



W wydaniu
wrześniowym m.in.:



PROJEKTY dla elektroników

- ➡ Generator funkcyjny Wi-Fi DDS, część 1
- ➡ Stylocclone – instrument muzyczny
- ➡ Cyfrowy terminal video z Raspberry Pi Pico, część 2

DIY dla wszystkich

- ➡ Jak używać darmowych rozwiązań opartych na sztucznej inteligencji
- ➡ Nocne światło alarmowe z Arduino

TUTORIALE

- ➡ Nauka elektroniki z modułem ESP32, część 2. Wejścia i wyjścia cyfrowe
- ➡ Kick Start część 12: Twój pierwszy oscyloskop
- ➡ Chirurgia obwodowa: Rezystancja sterowana elektronicznie, część 2
- ➡ Edukacja w EdW dla szkół i uczelni. Wykład 44: Płyta CD audio

JUNIOR

- ➡ Dwudzieste szóste spotkanie z najmłodszymi pasjonatami elektroniki





NIE PRZEOCZ

Nowe podzespoły	6
-----------------------	---

TEMAT NUMERU

Fotowoltaika balkonowa. Mały magazyn energii przy elektrowni 800 W – technika, bilans strat i polskie realia	14
--	----

AI W ELEKTRONICE

AI w pracowni elektronika konstruktora (4)	28
Kod z autopilota. Programowanie mikrokontrolerów wspomagane przez AI	32

Z PRAKTYKI EKSPERTÓW

Profesor Christophe Basso.	
Stabilizacja pętli trójfazowego prostownika Vienna (2)	53
Ochrona przepięciowa układów (1)	62

W PRACOWNI

Czujniki wilgotności i cieczy dla domu i ogrodu	72
Druk 3D w służbie elektroniki (11)	77

KITY AVT

Włącznik pomocniczy	80
---------------------------	----

PREZENTACJE

Obudowa to część układu. Jak uniknąć problemów, których nie widać na schemacie?	84
Projekt płytki na zlecenie	104

REPETYTORIUM

Obudowy w elektronice profesjonalnej	87
Zaawansowane funkcje oscyloskopów cyfrowych dla niezaawansowanych (1)	100
PCB w praktyce, część 1. Projektowanie PCB	110

KURS

Laboratorium pomiarowe dla konstruktora. Dwanaście rozdziałów – od oscyloskopu 12-bit po własny park pomiarowy	126
Oscyloskop 12-bit w praktyce konstruktora: ile bitów naprawdę trafia na ekran:	
Rozdział 1. 12 bitów to nie zawsze 12 bitów – co rozdzielczość i tryb High-Res dają konstruktorowi, a czego nie załatwią	127
Rozdział 2. Wyzwalanie i pomiary automatyczne: ukryta inteligencja oscyloskopu	132
Kurs RFID z PN532:	
Rozdział 5. MIFARE Ultralight i EV1: pamięć bez Crypto1	138
Rozdział 6. NTAG216: mirror UID i licznika – dynamiczne tagi bez zapisu	143

FELIETON

Czego producenci wielkich modeli językowych nie mówią inwestorom? Czyli kto i ile straci na nadchodzącym krachu?	149
--	-----

Prenumerata	2
Od wydawcy	3

Nowe podzespoły

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



Sterownik diod elektroluminescencyjnych NCR320U

NCR320U to sterownik pozwalający na zasilanie diod elektroluminescencyjnych (LED) stabilizowanym prądem wyjściowym o wartości: od 10 do 250 mA. Rozwiązanie pracuje poprawnie przy napięciu zasilania wynoszącym co najmniej 1,4 V, a maksymalne napięcie zasilające równa się 16 V. Wartość prądu wyjściowego dobierana jest z pomocą zewnętrznego rezystora dołączonego do wejścia REXT, umożliwiając łatwe i precyzyjne dopasowanie

parametrów pracy do wymagań konkretnej aplikacji. Zmiany prądu w funkcji napięcia są małe i nie przekraczają 1%/V, co gwarantuje absolutną stabilność pracy podzespołu. Dodatkowo temperaturowy współczynnik prądowy wynosi $-0,27\%/K$, a maksymalna, dopuszczalna moc strat to 750 mW, przy temperaturze otoczenia $25^{\circ}C$.

Układ zapewnia stabilizację prądu w całym, dopuszczalnym zakresie napięcia zasilania, dzięki czemu strumienie świetlne diod LED pozostają praktycznie stałe, niezależnie od wahań napięcia, czy zmian temperatury otoczenia. W trakcie pracy sterownik NCR320U pobiera nie więcej niż 1,2 mA prądu, co sprzyja jego efektywnej integracji z innymi podzespołami elektronicznymi. Przyjęte w układzie rozwiązania pozwalają na bezpośrednie sterowanie diodami LED, bez konieczności użycia dodatkowych, złożonych obwodów. Co więcej, możliwe jest równoległe łączenie kilku sterowników ze sobą, dzięki czemu komponent jest skalowalny i równocześnie elastyczny, jeżeli chodzi o zastosowania.

Sterownik NCR320U znajduje zastosowanie wszędzie tam, gdzie kluczowe jest zapewnienie stabilnego zasilania diod LED. Szczególnie często wykorzystywany jest w sektorze motoryzacyjnym, m.in. w systemach oświetlenia wnętrza pojazdów. Doskonale sprawdza się on również przy podświetlaniu kontrolki na deskach rozdzielczych i można go zastosować celem oświetlenia tablic rejestracyjnych znajdujących się z tyłu samochodów.

www.nexperia.com



Układ scalony ST25R210 do odczytu znaczników NFC/RFID

ST25R210 to scalony czytnik znaczników NFC i RFID, przeznaczony do komunikacji bezstykowej. Podzespół integruje w sobie tor nadawczo-odbiorczy, który odpowiada za realizację transmisji ze znacznikami, a także za generację pola elektromagnetycznego, niezbędnego do ich zasilania. Jest to rozwiązanie wspierające standardy: ISO/IEC 14443 oraz ISO/IEC 15693, współpracujące z szeroką klasą aplikacji, zapewniające kompatybilność z najpopularniejszymi typami kart bezstykowych, stosowanych m.in. w systemach identyfikacji, kontroli dostępu oraz urządzeniach IoT, zarówno komercyjnych, jak i przemysłowych.

Komunikacja z urządzeniami zewnętrznymi odbywa się przy pomocy interfejsu szeregowego, który pozwala na: konfigurację rejestrów, inicjalizację transmisji lub odbiór danych zgodnie z wybranym standardem. Wykrywanie znaczników bazuje na analizie zmian amplitudy lub fazy sygnału, a parametry toru nadawczo-odbiorczego mogą zostać zdefiniowane programowo. Dodatkowo oferowane są energooszczędne tryby pracy, w których aktywność czytnika zostaje obniżona do minimum. Okresowe załączanie czytnika umożliwia identyfikację znaczników przy mniejszym poborze mocy. Zasilanie z pojedynczego źródła oraz zintegrowane stabilizatory napięcia gwarantują niezawodną pracę, przy dostępności mechanizmów nadzorczych, które monitorują każde, niepożądane działanie i generują w odpowiedzi stosowne przerwy.

Czytnik ST25R210 umożliwia automatyczną realizację sekwencji transmisji zgodnych z wymaganiami zadanych protokołów. Dzięki temu układ z powodzeniem znajduje zastosowanie szczególnie w przemyśle, wyróżniając się niewrażliwością na zakłócenia elektromagnetyczne (EMI). Czyni to z niego efektywne, niezawodne rozwiązanie nawet dla aplikacji działających w trudnych środowiskach, gdzie stabilność komunikacji ma kluczowe znaczenie dla poprawnego funkcjonowania systemów w nich działających.

www.st.com



Regeneracyjne źródło zasilania RZ-X2-100K

RZ-X2-100K to dwukierunkowe źródło zasilania, zaprojektowane do testów aplikacji energoelektronicznych. Urządzenie umożliwia zarówno dostarczanie, jak i odbiór energii, przez co określane jest jako regeneracyjne. Szeroki zakres regulacji napięcia i prądu pozwala uzyskać moc wyjściową równą 400 kW, co czyni ze źródła produkt do badań m.in. falowników, systemów magazynowania energii i napędów w pojazdach elektrycznych.

Rozwiązanie znajduje zastosowania w kluczowych rozwiązaniach wysokonapięciowych. Kilka źródeł RZ-X2-100K może zostać ze sobą połączonych, co pozwala zwiększyć całkowitą moc systemu nawet do 2 MW. Taka skalowalność gwarantuje pełną elastyczność oraz umożliwia dostosowanie infrastruktury do przyszłych potrzeb rozwojowych.

Jedną z kluczowych funkcji źródła RZ-X2-100K jest emulacja baterii. Pozwala ona odwzorować charakterystyki napięciowo-prądowe, odpowiadające rzeczywistym warunkom pracy magazynów energii. Dzięki temu, możliwe staje się prowadzenie zaawansowanych testów szerokiej klasy rozwiązań. Dodatkowym atutem jest dostatecznie krótki czas odpowiedzi (poniżej 10 ms), co przekłada się na szybką i precyzyjną reakcję na zmiany obciążenia.

Przełączanie między trybami pracy odbywa się płynnie, bez przerw, eliminując niepożądane stany przejściowe i ograniczając „przeregulowania” prądu. Urządzenie wyposażono również w rozbudowany system zabezpieczeń, obejmujący w praktyce: monitoring napięcia, prądu i temperatury oraz wyłącznik różnicowoprądowy i funkcję awaryjnego zatrzymania. W sytuacjach zagrożenia gwarantuje to błyskawiczne odłączenie zasilania. Intuicyjny panel dotykowy umożliwia przy tym niezłożoną konfigurację parametrów oraz nieprzerwane śledzenie procesu testowego aplikacji.

Wysoka sprawność energetyczna źródła oraz niski poziom generowanego przez nie hałasu pozwalają na pracę w wymagających środowiskach. Co więcej, obsługa standardów komunikacji CAN i CAN FD upraszcza integrację

urządzenia z systemami nadrzędnymi. W ten sposób RZ-X2-100K stanowi niezawodne i elastyczne rozwiązanie dla nowoczesnych aplikacji, łącząc wydajność z łatwością integracji oraz komfortem użytkowania, umożliwiając równocześnie wierne odwzorowanie warunków pracy systemów elektromobilnych.

www.takasago-ss.co.jp



Sprzęgacz antenowy GC-130

Sprzęgacz antenowy GC-130 pozwala na dystrybucję sygnału nawigacji satelitarnej GPS do wielu odbiorników, przydając się tam, gdzie bezpośredni odbiór sygnału jest niemożliwy lub utrudniony. Rozwiązanie współpracuje z zewnętrznymi antenami aktywnymi, podłączanymi za pomocą kabla koncentrycznego, których średnice muszą nie przekraczać 14". Wejście antenowe sprzęgacza stanowi żeńskie gniazdo TNC. Parametry techniczne GC-130 obejmują straty wtrąceniowe na poziomie $8,1 \pm 0,5$ dB oraz straty odbiciowe typowo 12 dB. Dodatkowym atutem jest czerwony kolor obudowy, który pozwala na szybkie odnalezienie podzespołu w miejscu instalacji. Obudowa zapewnia z jednej strony skuteczne ekranowanie elektromagnetyczne, a z drugiej odporność mechaniczną, dzięki czemu produkt sprawdza się wszędzie, nawet w wymagających środowiskach pracy.

Instalacja sprzęgacza jest nadzwyczaj prosta – wyraźnie oznaczone wejście antenowe eliminuje ryzyko pomyłki

REKLAMA

BORNICO

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

✉ bornico@bornico.com.pl

🌐 www.bornico.com.pl

☎ +48 517 312 709 | +48 517 312 419

podczas podłączania. Jest to urządzenie o wymiarach: 63,25×60,2×33,27 cm, które cechuje stosunkowo nieduża waga (3,6 kg), przekładająca się zarówno na wygodę montażu, jak i niezłożony transport GC-130. Co kluczowe,

sprzęgacz firmy VIAVI Solutions może pracować w bezpośrednim sąsiedztwie innych anten (on też jest anteną), co potwierdza parametr separacji antenowej (antenna isolation) równy minimum 20 dB. W zestawie znajduje się kabel antenowy o długości około 6,1 m razem z walizką transportową. Rozwiązanie znajduje zastosowanie m.in. na stanowiskach testowych odbiorników GPS. Przydaje się ono także w hangarach lotniczych, w których sygnał satelitalny jest tłumiony oraz w terenie otwartym, gdzie dostęp do niego jest ograniczony i wymusza stosowanie dedykowanych podzespołów zapewniających stabilne, kontrolowane i powtarzalne warunki odbioru sygnału bez względu na otoczenie i warunki propagacji w nim obowiązujące.

www.viavisolutions.com



Termistorowe czujniki temperatury z serii NTCRP

Seria NTCRP obejmuje czujniki temperatury przewidziane do pracy w szczególnie wymagających środowiskach przemysłowych. Dzięki zastosowaniu precyzyjnego elementu pomiarowego, który stanowi termistor o ujemnym

współczynniku temperaturowym (NTC), istnieje możliwość przeprowadzania pomiarów w zakresie wartości: od -40 do 200°C, przy zachowaniu niskiego, a wręcz pomijalnego błędu. Termistor zamknięty jest w hermetycznej obudowie zabezpieczonej żywicą PPS (opartej na polisiarczku fenylenu). Takie rozwiązanie gwarantuje wysoką stabilność parametrów roboczych, chroniąc dostatecznie przed wpływem szkodliwych czynników zewnętrznych. Dodatkowo dostępne warianty konstrukcyjne pozwalają na łatwy montaż czujników zarówno na powierzchniach płaskich, jak i na przewodach o przekroju okrągłym, co ułatwia dopasowanie podzespołów do różnych aplikacji i konfiguracji instalacyjnych z nimi związanych.

Charakterystyczne dla czujników NTCRP są wyprowadzenia o długości 15...30 cm, odporne na substancje

oleiste i wysoką temperaturę otoczenia. Z najważniejszych parametrów do wymienienia są m.in.: rezystancja nominalna 3,3 kΩ, jej tolerancja ±5%, stała B₀/100 równa 3970 K oraz stała rozpraszania wynosząca 1,9 mW/°C, które razem pozwalają na szybkie i sprawne dopasowanie sensorów do wymogów aplikacji, przy jednoczesnym ograniczeniu stopnia nagrzewania się podzespołów. Jest to zagwarantowane obok możliwości pracy w miejscach wilgotnych i przy izolacji, którą w szczególności opisują: rezystancja powyżej 100 MΩ oraz napięcie przebicia 1,5 kV. Dzięki specjalnie przemyślanej konstrukcji, termistorowe czujniki temperatury z serii NTCRP wyróżniają się względnie niskim czasem reakcji. Umożliwia to ich zastosowania w szerokiej gamie zastosowań, wśród których występują: napędy w pojazdach elektrycznych, magazyny energii i układy energoelektroniczne, a także systemy zarządzania bateriami (BMS), falowniki i zasilacze, w których nacisk kładziony jest na efektywność pracy oraz sprawność, z czego ostatnie musi być wysokie, jeśli nie najwyższe.

www.tdk.com

Czujnik magnetyczny MRMS166R

MRMS166R to czujnik, którego zasada działania oparta jest na zjawisku magnetooporu anizotropowego (AMR), cechowany rezystancją zmieniającą się pod wpływem zewnętrznego pola magnetycznego. Dzięki wykorzystaniu zjawiska, podzespół może z łatwością spełniać rolę bezstykowego przełącznika. Brak elementów mechanicznych w elemencie eliminuje zużycie typowe dla klasycznych rozwiązań, co bezpośrednio przekłada się na zwiększoną trwałość komponentu oraz pełną niezawodność systemów, w których MRMS166R znajduje zastosowanie.



Podczas działania sensor firmy Murata Manufacturing wyróżnia się poborem prądu mniejszym od 20 nA. Dzięki tak niskiej konsumpcji energii, podzespół z powodzeniem może być używany w aplikacjach zasilanych z baterii, pracując nawet do dwóch lat, bez wymiany źródła zasilania. W celu zapewnienia prawidłowego funkcjonowania MRMS166R wymaga napięcia z przedziału: 1,2...3,6 V DC, przy wartości znamionowej, która wynosi 1,5 V DC. Równocześnie czujnik oferuje niezawodną i stabilną detekcję pola magnetycznego o indukcji równej od 0,5 do 4 mT, zachowując właściwości użytkowe również w zmiennych warunkach środowiskowych, operując poprawnie przy temperaturze otoczenia: od -10 do 60°C.

MRMS166R ma bardzo kompaktowe wymiary (1×1×0,4 mm), dzięki czemu można go łatwo zastosować

nawet w niedużych urządzeniach. Jest to element zaprojektowany z myślą o precyzyjnym wykrywaniu położenia oraz realizacji funkcji przełączających, przy obniżonym zużyciu energii. Najlepiej sprawdza się on w rozwiązaniach, w których decyduje oszczędność miejsca i wysoka efektywność energetyczna, takich jak: urządzenie przenośne, elektronika ubieralna, czy systemy IoT. Jednocześnie zapewniona jest pełna niezawodność działania, co ma duże znaczenie w przypadku eksploatacji – szczególnie w aplikacjach wymagających stabilnej i bezobsługowej pracy przez długi czas.

www.murata.com

Układ ZED-X20P-01B do odbioru sygnałów globalnych systemów nawigacji satelitarnej

ZED-X20P-01B to zaawansowany układ firmy u-blox kompatybilny z wieloma globalnymi systemami nawigacji satelitarnej (m.in. GPS i Galileo), umożliwiający pozycjonowanie z dokładnością do pojedynczych centymetrów. To rozwiązanie spełnia w praktyce rolę wysoce precyzyjnego, stabilnego źródła odniesienia, nawet w miejscach, w których nieprzerwanie występują zakłócenia sygnałów – przypadkowe bądź celowe, wpisujące się w działania dywersyjne obcych państw.



Zaimplementowana w podzespołe funkcja „moving base” pozwala na wyznaczanie położenia pomiędzy dwoma odbiornikami, co ma znaczenie w przypadku systemów przenośnych i autonomicznych. Element oferuje również

mechanizmy wykrywania zakłóceń odbieranego sygnału, dzięki czemu potrafi rozpoznać próby jego zagłuszenia (jamming), co jednoznacznie wykazano w ramach testów odporności na zakłócenia. Dodatkowo zachowana jest pełna zgodność mechaniczna i elektryczna z wcześniejszymi modułami serii ZED, co znacząco upraszcza proces modernizacji istniejących rozwiązań. Oznacza to zwłaszcza brak konieczności przeprojektowywania płytek PCB i dokonywania „głębszych” ingerencji w architekturę sprzętową systemów.

REKLAMA

Altium. Agile

**KUP NOWY
ALTIUM AGILE TEAMS
na 12 miesięcy,**

**ODBIERZ dodatkowe
2 miesiące GRATIS!**

Dowiedz się więcej na www.ccontrols.pl

COMPUTER
CONTROLS

Altium Agile Teams to zintegrowana platforma dla zaawansowanej współpracy multidyscyplinarnej.

www.ccontrols.pl
Oficjalny dystrybutor Altium

Dzięki powyższym właściwościom, układ ZED-X20P-01B znajduje zastosowanie w systemach bezzałogowych, robotyce, automatyce i rolnictwie precyzyjnym. Element sprawdza się także w aplikacjach pomiarowych, które operują w środowiskach o ograniczonej infrastrukturze, zapewniając stabilne i dokładne pozycjonowanie mimo utrudnionego odbioru sygnałów radiowych oraz zmiennych warunków ich propagacji, które są tego przyczyną.

www.u-blox.com



Czujnik parametrów środowiskowych SEN62

SEN62 to czujnik, który służy do kompleksowego monitorowania jakości powietrza, umożliwiając równoczesny pomiar kilku zasadniczych parametrów środowiskowych. Rozwiązanie rejestruje: stężenia pyłów zawieszonych (PM1, PM2.5 i PM10), poziom lotnych związków organicznych (VOC), a także, co ważne, wilgotność względną i temperaturę otoczenia.

Pomiar stężenia pyłów realizowany jest metodą optyczną, opartą na zjawisku rozpraszania światła. Wbudowany w czujnik wentylator zapewnia wymuszony przepływ powietrza przez komorę pomiarową SEN62, co w efekcie przekłada się na stabilne, powtarzalne wyniki pomiarów. Dodatkowo stosowane są mechanizmy kompensacyjne, zdecydowanie ograniczające dryft wskazań w warunkach długotrwałej eksploatacji podzespołu.

Detekcja lotnych związków organicznych odbywa się z wykorzystaniem dedykowanego elementu pomiarowego, charakteryzującego się krótkim czasem reakcji. W podobny sposób realizowane są również pomiary temperatury i wilgotności względnej, dzięki czemu wszystkie istotne dane zbierane są w sposób spójny i zsynchronizowany.

Do komunikacji z systemami nadrzędnymi wykorzystuje się interfejs cyfrowy, który umożliwia łatwą integrację czujnika SEN62 z różnorodnymi aplikacjami – w tym z programowalnymi sterownikami logicznymi (PLC) oraz systemami zarządzania budynkiem (BMS). Jednocześnie niebagatelnym atutem rozwiązania firmy Sensirion jest fabryczna kalibracja, dzięki której na etapie instalacji nie ma konieczności dokonywania dodatkowych

regulacji, co upraszcza i przyspiesza uruchomienie sensora. Dodatkowo w czujniku oferowane są funkcje autodiagnostyki umożliwiające bieżące monitorowanie poprawności działania podzespołu oraz wykrywanie ewentualnej degradacji elementów pomiarowych w nim zawartych.

Czujnik SEN62 wymaga zasilania stabilizowanym napięciem, co warunkuje jego niezawodną oraz powtarzalną pracę. Wbudowane w element układy regulacyjne odpowiadają za utrzymanie właściwych parametrów działania produktu, a dzięki wytrzymałej konstrukcji mechanicznej rozwiązanie może być montowane wszędzie, nawet w kanałach wentylacyjnych. W jednym komponencie integruje się wiele kluczowych funkcjonalności, co znacząco upraszcza budowę systemów monitorowania jakości powietrza, prowadząc z jednej strony do obniżenia kosztów wdrażania, a z drugiej do zwiększenia niezawodności aplikacji, w którym SEN62 znajduje swoje zastosowanie w sposób szeroki, jeśli nie najszerszy.

<https://sensirion.com>



Pastylkowa bateria litowo-manganowa CR2450J

Bateria CR2450J firmy Jauch Quartz została stworzona przede wszystkim z myślą o elektronicznych etykietach cenowych, choć z powodzeniem sprawdza się także w innych energooszczędnych rozwiązaniach. Jest to ogniwo litowo-manganowe (Li/MnO₂) w kształcie pastylki o średnicy 24,5 mm i wysokości 5 mm, które oferuje napięcie nominalne 3 V oraz pojemność maksymalną 620 mAh. Co istotne, spadek pojemności po kilku latach użytku jest stosunkowo niewielki i zazwyczaj mieści się w przedziale 10...20%.

Na szczególną uwagę zasługuje niski poziom samorozładowania, który przekłada się na długą żywotność baterii (np. w systemach podtrzymania pamięci). Co więcej, zaletą jest także wysoka stabilność napięcia, która oznacza niezawodną pracę różnego typu urządzeń elektronicznych. Prawidłowe działanie ogniwa CR2450J obliczone jest przy tym na zakres temperatury: od -30 do 85°C, co pozwala na wykorzystanie elementu w mniej wymagających środowiskach. Dodatkowo należy pamiętać,

że najwyższy prąd ciągły wynosi zaledwie kilka miliamperów – jego przekroczenie może skutkować spadkiem napięcia oraz przyspieszonym zużyciem baterii.

Ogniwo spełnia wymagania dyrektyw środowiskowych RoHS oraz REACH i nie zawiera ołowiu, co potwierdza jego zgodność z normami ograniczającymi stosowanie substancji niebezpiecznych w sprzęcie elektrycznym oraz elektronicznym. Dzięki temu, bateria może być bezpiecznie stosowana w nowoczesnych urządzeniach, a atutem jest również dostępność kilku wariantów montażowych. Element dostarczany jest w wersjach przystosowanych do montażu przewlekane lub powierzchniowego (SMT) oraz wyposażonych w przewody, co znacząco ułatwia jego integrację z dowolnymi rozwiązaniami – niezależnie od specyfiki projektu, wymagań konstrukcyjnych, czy technologii produkcji, którą się wykorzystuje.

www.jauch.com



Pełnospektralne diody elektroluminescencyjne SunLike

Określenie „pełnospektralne” odnosi się do źródeł światła, których widmo jest możliwie najbardziej zbliżone do światła naturalnego, czyli słonecznego. W diodach elektroluminescencyjnych serii SunLike ma miejsce zupełnie inne podejście niż w klasycznych diodach LED – zamiast emitera światła niebieskiego stosuje się emiter światła fioletowego wraz ze specjalnie dobranym luminoforem reagującym na nową długość fali. Dzięki temu, uzyskuje się bardziej równomierny rozkład widma w całym zakresie światła widzialnego – bez wyraźnych „pików”, które obecne są w przypadku diod standardowych. Generowane światło staje się bliższe naturalnemu, a współczynnik oddawania barw (CRI) osiąga bardzo wysokie, często najwyższe, maksymalne wartości.

Diody serii wyróżniają się dużą elastycznością zastosowań, współpracując z licznymi układami sterowania oraz przetwornicami pozwalającymi na bezpośrednią integrację podzespołów w systemach oświetleniowych. Dzięki temu, rozwiązania sprawdzają się wprost w instalacjach zarówno profesjonalnych, jak i bardziej standardowych,

dla których najważniejsze są: niezawodna praca oraz prostota wdrożenia.

Na tle pozostałych diod LED seria SunLike wyróżnia się przede wszystkim znacznym ograniczeniem zmęczenia oczu. Wyniki badań wskazują także na zahamowanie postępów krótkowzroczności przy długotrwałej ekspozycji na światło generowane przez elementy serii. Oznacza to, że zdecydowanie poprawiają one komfort widzenia, co ma kluczowe znaczenie podczas wielogodzinnej pracy wymagającej skupienia. Podzespoły najlepiej sprawdzają się w miejscach, gdzie wzrok jest intensywnie eksploatowany – takich jak: szkoły, uczelnie czy inne środowiska pracy umysłowej. Szczególnie korzystny wpływ diod ma miejsce w pomieszczeniach z ograniczonym dostępem do światła dziennego, ponieważ elementy ograniczają tam skutecznie negatywny wpływ oświetlenia sztucznego na organizm, nie wymagając przy tym zmian w instalacji elektrycznej, ani zapewniania specjalnych warunków, co czyni je produktami wyjątkowo prostymi w implementacji i atrakcyjnymi z punktu widzenia użytkowników końcowych.

www.seoulsemicon.com

Czujnik optyczny SHF-7061

Czujnik SFH 7061 przeznaczony jest do wykrywania zarówno nacisków, jak i obecności obiektów znajdujących się w bardzo małej odległości. Działanie sensora opiera się bezpośrednio na współpracy elementów, którymi są: emiter promieniowania podczerwonego (IR) i fotodiody odbiorcza. Oba komponenty zawarte są w kompaktowej obudowie, określonej przez wymiary: 1,8×1×0,5 mm, pozwalającej na stosowanie podzespołu w miniaturowych rozwiązaniach elektronicznych.



Emiter podczerwieni pracuje przy długości fali 950 nm, charakteryzując się szerokością spektralną równą 29 nm.

REKLAMA

PRODUCENT
**ELEMENTÓW
INDUKCYJNYCH**

RY

www.feryster.pl

Q FERYSTER

Napięcie przewodzenia VF związane z emiterem mieści się w przedziale: od 1,3 do 2,9 V. Równocześnie należy zwrócić szczególną uwagę na napięcie wsteczne VR: jego przekroczenie ponad wartość nominalną 5 V prowadzi do nieodwracalnego uszkodzenia SHF-7061.

Zastosowana fotodioda reaguje na promieniowanie o długości fali z zakresu: 820...1100 nm, osiągając najwyższą czułość w pobliżu 910 nm. Generowany prąd fotoelektryczny IP osiąga wartość do około 90...100 nA. Natomiast prąd ciemny, uzależniony od napięcia VR, jest mniejszy od 0,4 nA. Dzięki tak korzystnym parametrom, staje się możliwe znaczące zmniejszenie szumów, co przekłada się na wysoce precyzyjną detekcję nawet już niewielkich zmian natężenia promieniowania.

SHF-7061 przystosowany jest do pracy w temperaturze: od -40 do 85°C, co pozwala na wykorzystanie czujnika w wymagających środowiskach. Dodatkowo układ spełnia standard ANSI/ESDA/JEDEC JS-001 potwierdzający odporność na wyładowania elektrostatyczne (ESD), mające miejsce zarówno w trakcie montażu, jak i codziennej eksploatacji urządzeń.

W zastosowaniach praktycznych rozwiązanie współpracuje z analogowymi torami przetwarzania sygnału. Nawet najmniejsze zmiany odległości bądź właściwości odbiciowych powierzchni warunkują wartość prądu IP, a sygnał z nim związany należy poddać dalszemu przetwarzaniu obejmującemu: wzmacnianie, filtrowanie i próbkowanie, celem konwersji do postaci cyfrowej i dalszej analizy zwłaszcza w systemach mikroprocesorowych lub układach sterowania.

<https://ams-osram.com>



Wzmacniacz światłowodowy FX-250

FX-250 to rozwiązanie opracowane przez firmę Panasonic z myślą o wygodnej obsłudze oraz niezawodnej pracy w nowoczesnych aplikacjach. Jest tu mowa o podzespolu, które spełnia funkcję wzmacniacza światłowodowego, oferującym m.in. czytelny, wysokokontrastowy wyświetlacz OLED gwarantujący szybkie – zajmujące niepełną kilka sekund, sprawdzenie stanu pracy oraz bieżące ustawienia urządzenia.

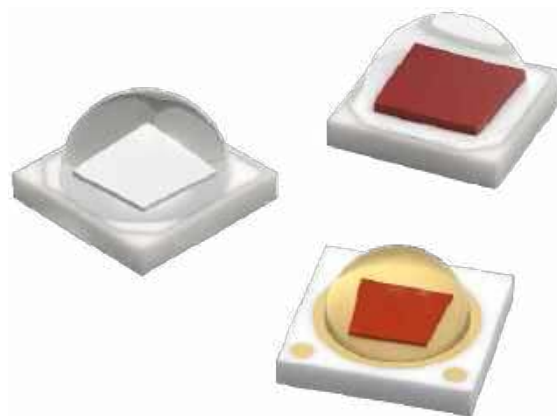
Kluczowym atutem FX-250 jest inteligentna, dwuetapowa funkcja automatycznego uczenia progu przełączania. Po dwukrotnym uruchomieniu procedury element samodzielnie dobiera optymalne parametry pracy, co pozwala utrzymać wysoką skuteczność działania nawet w trudnych warunkach. Wraz z nią wyjątkowo krótki czas reakcji wynoszący 35 µs sprawia, że urządzenie sprawdza się bezpośrednio na każdej linii produkcyjnej, równocześnie ograniczając wydatnie ryzyko popełnienia zasadniczych błędów konfiguracyjnych.

Kompaktowa konstrukcja podzespołu została prześlana pod kątem m.in. instalacji o ograniczonej przestrzeni montażowej, typowych dla współczesnych maszyn. Minimalny promień otwarcia pokrywy, wynoszący 3 cm, w istocie rzeczy przekłada się na bezpieczny dostęp serwisowy w trudno dostępnych miejscach, bez konieczności demontażu sąsiednich elementów. Na dodatek urządzenie wyposażone zostało w funkcje oszczędzania energii, które samoczynnie ograniczają pobór mocy poprzez przyciemnienie lub wyłączenie wyświetlacza i wskaźników po okresie bezczynności. Ma to szczególne znaczenie w systemach działających w trybie ciągłym, kiedy pobór mocy wpływa bezpośrednio na koszty eksploatacji.

Wzmacniacz FX-250 stanowi ekonomiczne rozwiązanie w kontekście całkowitego kosztu użytkowania (TCO).

Jego przewaga wynika przede wszystkim z zastosowania wysokiej jakości komponentów, które przekładają się na dłuższą żywotność urządzenia, rzadsze przestoje oraz niższe koszty serwisowania. W dłuższym horyzoncie oznacza to stabilne parametry pracy, a za zaletę uważana jest także prosta integracja z dostępną infrastrukturą, bez konieczności wprowadzania zmian oraz przy minimalnym zaangażowaniu zasobów tak technicznych, jak i organizacyjnych (bez względu na sytuację).

<https://industry.panasonic.eu>



Najnowsze diody elektroluminescencyjne WL-SMDC przeznaczone do upraw roślinnych

Seria diod WL-SMDC jest znana na rynku od minimum kilku lat i zyskała uznanie dzięki swojej niezawodności.

Z biegiem czasu linia podzespołów została rozwinęta i przystosowana do nowych wymagań, czego efektem jest rozszerzenie oferty o warianty przeznaczone do zastosowań w uprawach roślin. Nowe modele diod opracowano z myślą o potrzebach współczesnego ogrodnictwa, w którym znaczenie ma wszechstronne wspieranie procesów biologicznych. Komponenty przystosowane są do montażu powierzchniowego (SMT) i dostępne w standardowej obudowie 3535, co ułatwia ich integrację w szerokiej klasie systemów oświetleniowych.

Najważniejszym atutem nowych wariantów diod WL-SMDC jest optymalizacja ich parametrów świetlnych pod kątem efektywności fotosyntezy. Dzięki temu, rozwiązania lepiej odpowiadają na potrzeby roślin, wspierając ich rozwój w kontrolowanych warunkach uprawy. O jakich parametrach jest dokładnie mowa? Odpowiedź brzmi:

- fotosyntetyczna wydajność fotonowa (Photosynthetic Photon Efficacy – PPE) na poziomie 4,9 $\mu\text{mol/J}$,
- strumień fotonów fotosyntetycznych (Photosynthetic Photon Flux – PPF) sięgający 6,34 $\mu\text{mol/s}$.

Dostępne warianty emitują światło o długościach fali: 450, 660 lub 730 nm, co pozwala precyzyjnie kształtować widmo promieniowania w paśmie kluczowym dla fotosyntezy.

We wszystkich diodach zapewniona jest możliwość płynnej regulacji prądu przewodzenia, w celu precyzyjnego sterowania natężeniem światła. Co więcej, w środowisku symulacyjnym REDEXPERT udostępniono narzędzie obliczeniowe, umożliwiające dobór parametrów oświetlenia, w zależności od gatunku rośliny, fazy wzrostu oraz wymaganych parametrów jakościowych. Wysoka skuteczność energetyczna oraz zwiększona gęstość strumienia światła umożliwiają projektowanie modułowych układów oświetleniowych o zredukowanym poborze mocy oraz o zwiększonej powierzchni pokrycia.

Nowe diody znajdują zastosowanie przede wszystkim w nowoczesnych systemach upraw – takich jak: szklarnie, farmy pionowe lub instalacje kontenerowe. Co kluczowe, parametry optoelektroniczne podzespołów dobrano na podstawie licznych badań naukowych. Jednocześnie elementy pozwalają na precyzyjne kształtowanie warunków świetlnych w uprawach, umożliwiając precyzyjne dopasowanie widma światła do potrzeb fizjologicznych roślin, co wspiera ich wzrost, poprawia efektywność produkcji oraz ostatecznie zapewnia lepszą kontrolę, ograniczając przy tym zużycie energii oraz dodatkowo minimalizując koszty eksploatacji całego systemu.

www.we-online.com



Moduły mocy QSiC Dual3

Pierwotnie moduły QSiC Dual3 – o napięciu znamionowym 1200 V, powstały z myślą o budowie nowoczesnych przekształtników energoelektronicznych. W podzespołach, których producentem jest firma SemiQ, stosowane są tranzystory MOSFET wykonane z węgla krzemu (SiC), cechujące się niską rezystancją w stanie przewodzenia $R_{DS(on)}$ z przedziału: 1...2 m Ω . Takie wartości parametru przekładają się wprost na możliwie najwyższą sprawność pracy rozwiązań. W zależności od wariantów moduły osiągnęły gęstość mocy na poziomie 240 W/in³, zachowując przy tym kompaktową obudowę o wymiarach: 6,2×15,2 cm. Dodatkowym atutem jest efektywne odprowadzanie ciepła, które pozwala obywać się bez stosowania aktywnych układów chłodzenia.

Seria QSiC Dual3 została pomyślana jako atrakcyjna alternatywa dla modułów opartych na tranzystorach IGBT. Jej zastosowanie zwykle nie wymaga wprowadzania gruntownych zmian w konstrukcji urządzeń – w większości przypadków wystarczają drobne korekty z zakresu sterowania i torów mocy. Korzyścią jest również zmniejszona waga rozwiązań, a ponadto ich niewielkie gabaryty zdecydowanie ułatwiają integrację z pozostałymi elementami systemu.

Moduły serii znajdują zastosowania w szerokim wachlarzu aplikacji. Wykorzystywane są one m.in. w: napędach silników elektrycznych i przekształtnikach sieciowych obecnych w systemach magazynowania energii, a także w układach zasilania i chłodzenia centrów danych. Decydujące zalety produktów stają się zauważalne zwłaszcza w aplikacjach dużej mocy – takich jak np. agregaty chłodnicze. W tego rodzaju rozwiązaniach moduły QSiC Dual3 pozwalają obniżyć koszty eksploatacyjne, jednocześnie zwiększając niezawodność całego systemu. Dzięki temu, urządzenia mogą pracować stabilnie nawet w trudnych warunkach, ograniczając potrzebę częstych przeglądów i serwisowania w perspektywie miesięcy, wielu lat, a nawet dłużej.

<https://semi-q.com>



Fotowoltaika balkonowa

Mały magazyn energii przy elektrowni 800 W – technika, bilans strat i polskie realia

Elektrownia balkonowa to jeden z niewielu przypadków, w których urządzenie energetyczne rozeszło się po Europie szybciej, niż nadążyli za nim ustawodawcy: dwa moduły fotowoltaiczne na balkonie, mikrofalownik, wtyczka do gniazdka i limit 800 W, który stał się w praktyce ogólnoeuropejskim punktem odniesienia. Thomas Scherer, autor czasopisma Elektor, należy do entuzjastów tej konstrukcji od samego początku – zainstalował ją u siebie ponad pięć lat temu i od tego czasu regularnie o niej pisze. W numerze majowo-czerwcowym 2026 opisał kolejny krok: dołożenie do takiej instalacji małego magazynu energii klasy 2 kWh [1]. Rzecz w tym, że przełom, o którym pisze, nie nastąpił w ogniwach ani w przetwornicach, lecz w oprogramowaniu i w zewnętrznym pomiarze prądu. I w tym, że autor kupił ten magazyn z powodu, który w Polsce nie istnieje – u niego nadwyżka z balkonowej elektrowni przepada bezpowrotnie, u nas trafia na depozyt prosumencki. Czy zatem polskiemu użytkownikowi magazyn jest niepotrzebny? Niekoniecznie – tylko opłaca mu się z zupełnie innego powodu. Przyjrzyjmy się najpierw technice, a potem rachunkom w polskich realiach.

1. Skąd wziąć się magazyn przy balkonie

Sama elektrownia balkonowa jest układem trywialnym: jeden lub dwa moduły fotowoltaiczne o łącznej mocy 800...1000 Wp (watopik – szczytowa moc panelu fotowoltaicznego), mikrofalownik ograniczony do 800 W po stronie AC, przewód do gniazda w mieszkaniu. Energia wprowadzona do instalacji lokalu obniża wskazanie licznika, o ile w tej samej chwili jest

w mieszkaniu zużywana. Problem polega na tym, że dobowy profil produkcji i dobowy profil poboru różnią się w sposób systematyczny. Szczyt produkcji wypada w godzinach, w których w typowym mieszkaniu nikogo nie ma; szczyt poboru – wieczorem, gdy moduły już nie pracują. Efektem jest autokonsumpcja rzędu 30...50%, a więc połowa lub więcej wyprodukowanej energii trafia poza mieszkanie.

To, co dzieje się z tą nadwyżką, zależy od kraju i to jest wątek, do którego wrócimy w drugiej części artykułu. W Niemczech, o których pisze Scherer, przy elektrowni balkonowej nie zawiera się umowy o rozliczaniu energii wprowadzanej do sieci – nadwyżka jest po prostu oddawana operatorowi bez wynagrodzenia. Autor odnotowuje przy tym wyjątek, który dobrze pokazuje, jak bardzo rachunek zależy od szczegółu regulacyjnego: w kilku krajach, m.in. w Holandii i Luksemburgu, tolerowana jest praca elektrowni balkonowej ze starym licznikiem indukcyjnym bez blokady wstecznej. Tam licznik cofa się, sprawność takiego „magazynu” wynosi 100%, koszt zero i żaden akumulator nie ma sensu ekonomicznego.

Wszędzie tam, gdzie licznik ma blokadę, pojawia się pokusa dołożenia baterii. Pierwsza fala takich urządzeń, sprzed dwóch lat, wypadła jednak słabo. Magazyny były sterowane w sposób, który dziś nazwalibyśmy ślepym: użytkownik nastawiał moc oddawaną do instalacji, konfigurowalną często skokowo co ok. 100 W, i urządzenie trzymało tę wartość niezależnie od tego, co działo się w mieszkaniu. Skutki są łatwe do przewidzenia. Przy nastawie zbyt niskiej magazyn nie pokrywał bieżącego poboru i różnicę i tak trzeba było kupić w sieci. Przy nastawie zbyt wysokiej – i to jest przypadek gorszy – nadmiar wpływał przez licznik do sieci, czyli energia raz już zmagazynowana ze stratami była oddawana za darmo. Do tego dochodzi koszt, o którym łatwo zapomnieć: energia wyjęta z akumulatora nie jest darmowa, bo obciąża ją zużycie cykliczne ogniów i amortyzacja samego urządzenia.

Zmiana, którą opisuje Scherer, przyszła pod koniec 2025 roku i miała charakter programowy. Producenci nauczyli swoje magazyny korzystać z zewnętrznego pomiaru prądu w rozdzielnicy, a to otworzyło dwie funkcje, które dopiero czynią z takiego urządzenia sensowny element instalacji:

- **ładowanie nadwyżkowe** – magazyn pobiera energię wyłącznie wtedy, gdy produkcja modułów fotowoltaicznych przekracza bieżące zapotrzebowanie budynku; dopóki dom konsumuje wszystko, co produkują panele, bateria się nie ładuje;
- **zerowe wprowadzanie** (*zero feed-in*) – magazyn oddaje do instalacji dokładnie tyle mocy, ile w danej chwili jest potrzebne, i ani wata więcej.

Dopiero z tymi dwiema funkcjami magazyn przestaje być gadżetem, a staje się buforem, który przesuwa produkcję z południa na wieczór, nie tracąc jej po drodze do sieci. Scherer pisze, że to właśnie zadziałanie obu mechanizmów w klasie urządzeń 2 kWh – a więc tych realnie dobieranych do zestawu balkonowego, nie do instalacji dachowej – przekonało go do zakupu pod koniec grudnia 2025 roku.

Warto w tym miejscu zaznaczyć jedną rzecz. Ładowanie nadwyżkowe i zerowe wprowadzanie rozwiązują niemiecki problem – utraty nadwyżki. W polskim net-billingu nadwyżka nie jest tracona, więc te same funkcje pełnią u nas nieco inną rolę: nie ratują energii przed przepadnięciem, lecz zamieniają energię wycenianą tanio (jako wprowadzoną do sieci) na energię drogą (jako uniknięty zakup). To różnica ilościowa, nie jakościowa, ale ma bardzo konkretne przełożenie na okres zwrotu inwestycji – do czego dojdziemy w punkcie 6.

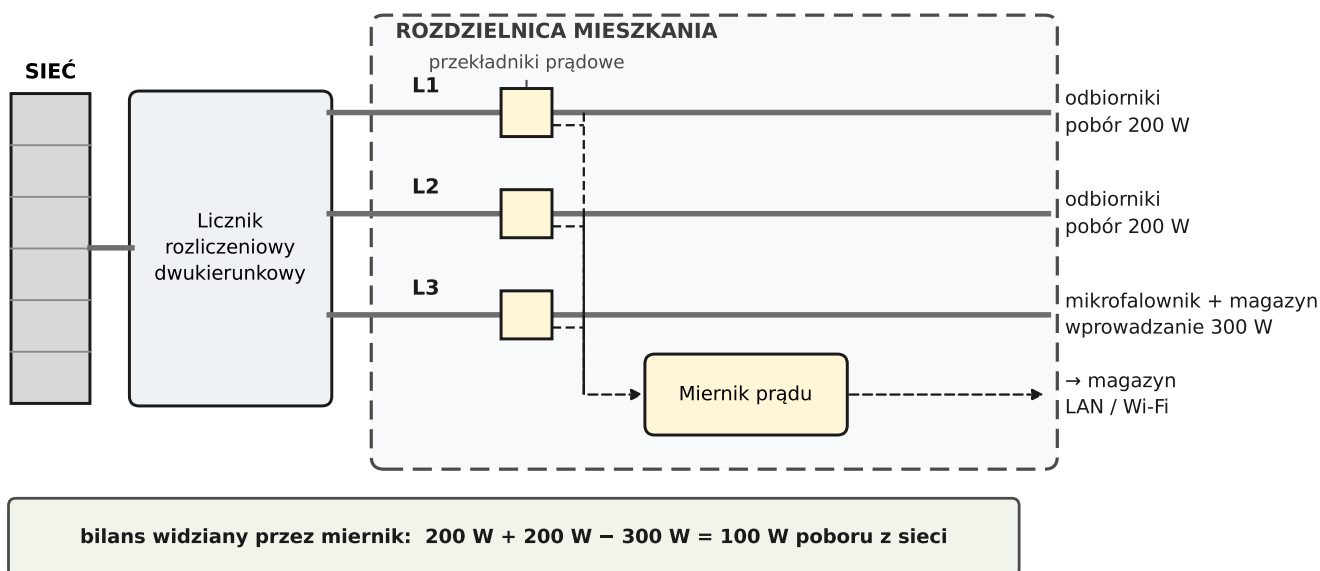
2. Warunek konieczny: pomiar prądu w rozdzielnicy

Zerowe wprowadzanie jest funkcją regulacyjną i jak każda regulacja wymaga sygnału sprzężenia zwrotnego. Magazyn musi wiedzieć, ile mocy przepływa przez złącze między instalacją budynku a siecią, i to wiedzieć na tyle szybko, by nadążyć za załączaniem odbiorników. Sam falownik tego nie wie – widzi wyłącznie własny prąd wyjściowy.

Stąd trzeci element zestawu, obok modułów fotowoltaicznych i magazynu energii: licznik prądu montowany w rozdzielnicy, w miejscu, w którym widać sumaryczny przepływ. Scherer odnotowuje, że rolę standardu de facto przejęły w tym segmencie urządzenia Shelly; sam zastosował Shelly Pro 3EM z przekładnikami o zakresie 120 A na fazę, w cenie 70..80 €. Zakres jest dla mieszkania mocno przewymiarowany, ale urządzenie jest dojrzałe i obsługiwane przez większość magazynów na rynku, co przy tej klasie sprzętu bywa argumentem ważniejszym od dopasowania zakresu.

Istotna jest zasada bilansowania. Miernik sumuje moce trzech faz algebraicznie, tak jak robi to licznik rozliczeniowy. Autor ilustruje to przykładem, który warto zapamiętać, bo bywa źródłem nieporozumień: przy poborze 200 W na fazie L1, 200 W na L2 i jednoczesnym wprowadzaniu 300 W na L3 miernik raportuje 100 W pobierane z sieci – i tę wartość magazyn próbuje skompensować, dopóki ma zapas energii. Rozkład obciążeń między fazami w mieszkaniu lub domu jest przy tym w praktyce przypadkowy, wynika z historycznego układu instalacji i nie da się go łatwo zmienić. Dla sterowania magazynem nie ma to jednak znaczenia – liczy się wyłącznie bilans.

Druga różnica względem „dużych” instalacji dotyczy warstwy komunikacyjnej. W klasycznym układzie prosumenckim falownik rozmawia z licznikiem po kablu, magistralą CAN lub RS485. Tutaj miernik komunikuje się z magazynem przez LAN lub Wi-Fi. Konsekwencje są dwie. Po pierwsze, magazyn musi być wpięty w domową sieć danych, co dla części użytkowników jest barierą wdrożeniową większą niż sam montaż mechaniczny. Po drugie – i to jest wniosek, który Scherer formułuje mimochodem,



Rysunek 1. Miejsce włączenia miernika prądu w rozdzielnicy i zasada bilansowania trójfazowego. Przekładniki obejmują przewody fazowe za zabezpieczeniem głównym; miernik przekazuje magazynowi wypadkową moc bilansową, a nie moce poszczególnych faz

a który dla czytelnika EP jest może najważniejszy – magazyn balkonowy jest urządzeniem klasy IoT ze wszystkimi tego konsekwencjami: zależnością od jakości firmware'u, od dostępności chmury producenta i od tego, czy producent zechce urządzenie wspierać za pięć lat. Autor odnotowuje zresztą, że jego instalacja zaczęła działać poprawnie dopiero po aktualizacji oprogramowania z 25 stycznia 2026 roku, a więc miesiąc po zakupie. To dobra ilustracja tego, jak młody jest ten segment rynku.

3. Trzy topologie

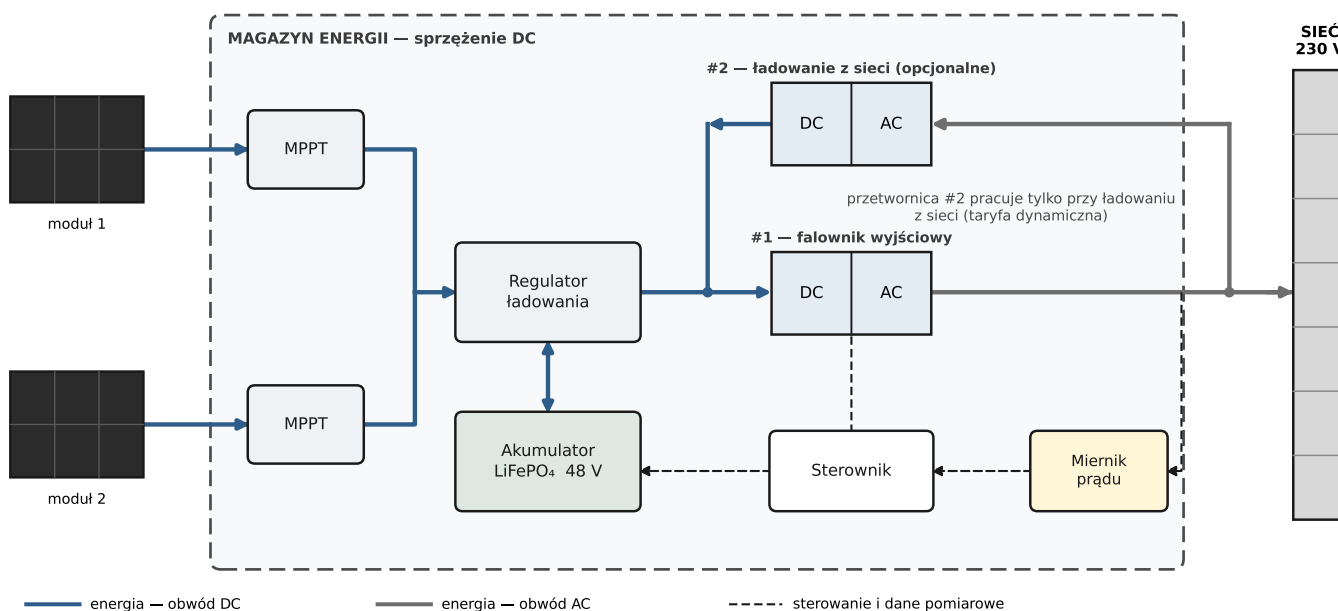
To jest sedno zagadnienia i punkt, w którym decyzja projektowa waży najwięcej. Pytanie brzmi: w którym miejscu toru energetycznego wpiąć akumulator – przed falownikiem, po stronie stałoprądowej, czy za nim, po stronie

sieciowej. Odpowiedź przesądza o sprawności, o dopuszczalnym miejscu montażu, o przekroju przewodów, a w Niemczech również o legalności układu.

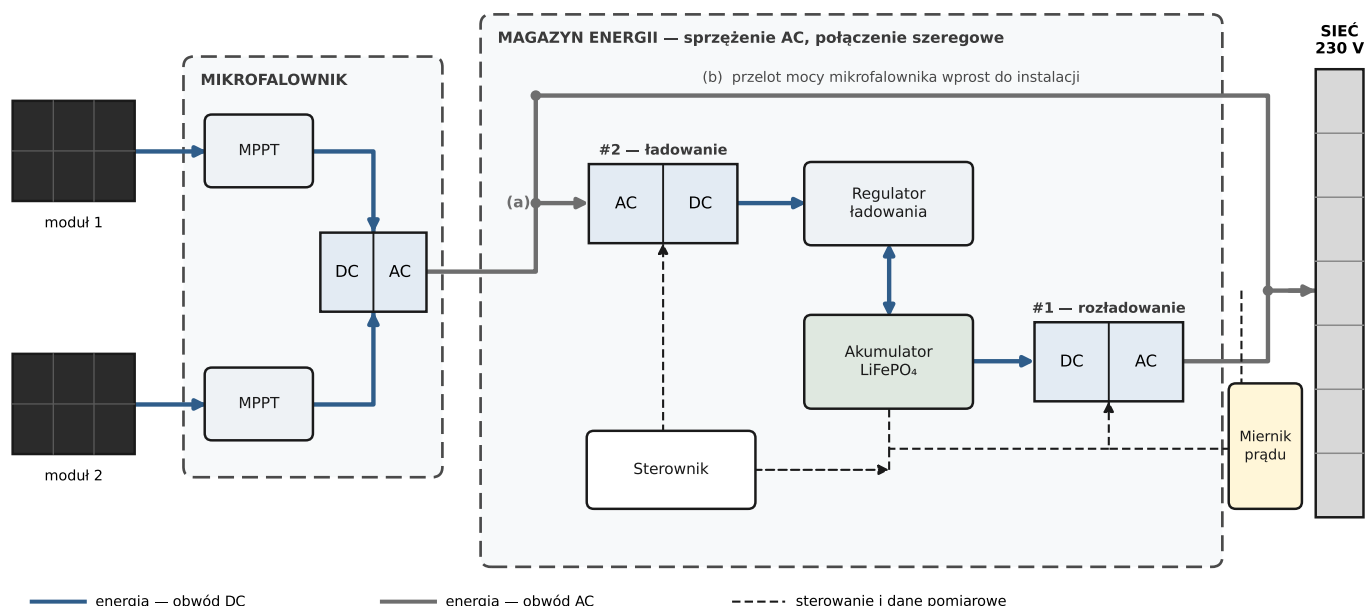
3.1. Sprzężenie DC

Moduły fotowoltaiczne przyłącza się bezpośrednio do magazynu, który zawiera regulator ładowania z układem MPPT. Zewnętrzny mikrofalownik jest w tej konfiguracji zbędny – jego funkcję pełni pojedyncza przetwornica DC/AC na wyjściu magazynu, podnosząca napięcie akumulatora (typowo 48 V) do napięcia sieci.

Zaletą jest uniknięcie podwójnej konwersji. W układzie ze sprzężeniem AC energia z modułów jest najpierw podnoszona z 30...45 V DC do 230 V AC, a następnie obniżana z powrotem do 48 V DC akumulatora – dwa przetwarzania,



Rysunek 2. Sprzężenie DC. Moduły przyłączone bezpośrednio do wejść MPPT magazynu, pojedyncza przetwornica DC/AC na wyjściu. Energia z modułów wchodzi do ogniwa przez jedno przetwarzanie



Rysunek 3. Sprzężenie AC szeregowe. Wyjście mikrofalownika wprowadzone do magazynu; widoczna podwójna konwersja DC/AC/DC po stronie wejściowej

dwie porcje strat, obie ponoszone jeszcze przed wejściem energii do ogniw. Tutaj tego nie ma: energia z modułów wchodzi do baterii przez jeden regulator.

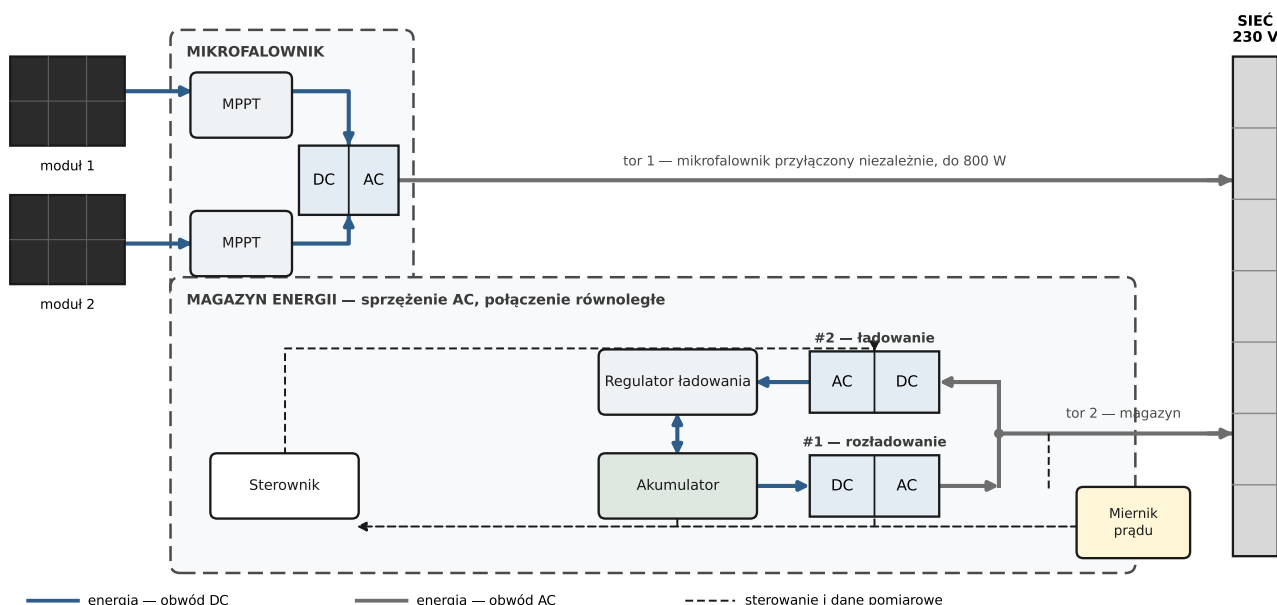
Wadą jest miejsce montażu. Moduły są na balkonie, więc magazyn musi być blisko nich, czyli na zewnątrz. Do tego dochodzi konieczność poprowadzenia sztywnych przewodów solarnych o zalecany przekroju 6 mm² na dłuższych odcinkach – kabla, którego nie da się poprowadzić przez uszczelkę okna ani ukryć pod listwą.

Wariantem rozszerzonym jest układ z dodatkową przetwornicą AC/DC na wejściu, pozwalającą ładować magazyn również z sieci. Ma to sens wyłącznie przy taryfie dynamicznej, gdy różnica cen między godzinami ładowania i rozładowania pokrywa straty przetwarzania

i amortyzację. Scherer sygnalizuje tę funkcję, ale świadomie zostawia ją na boku – jego artykuł dotyczy buforowania energii słonecznej, nie arbitrażu. Dla polskiego czytelnika akurat ten wątek okaże się kluczowy i wrócimy do niego w punkcie 7.

3.2. Sprzężenie AC szeregowe

Tu wykorzystuje się mikrofalownik, który w każdej elektrowni balkonowej i tak jest. Jego wyjście AC wprowadza się do magazynu, a magazyn łączy się z siecią. Wewnątrz magazynu energia przechodzi przez przetwornicę AC/DC do ogniw, a następnie przez przetwornicę DC/AC z powrotem do sieci – stąd wspomniana podwójna konwersja po stronie wejściowej.



Rysunek 4. Sprzężenie AC równoległe. Mikrofalownik i magazyn przyłączone do instalacji niezależnie, każdy własnym torem. Konfiguracja o najprostszym sterowaniu i najwyższej mocy chwilowej

Zaletą jest swoboda montażu. Magazyn może stać w mieszkaniu, a między mikrofalownikiem na balkonie a magazynem wewnątrz wystarczy zwykły przewód zasilający $3 \times 1,5 \text{ mm}^2$ zamiast czterech sztywnych przewodów solarnych 6 mm^2 . Dla instalacji w bloku to argument rozstrzygający.

3.3. Sprzężenie AC równoległe

Mikrofalownik i magazyn są przyłączone do instalacji niezależnie od siebie, każdy własnym przewodem. Wiele magazynów AC obsługuje obie konfiguracje – szeregową i równoległą – ale Scherer zaleca sprawdzenie tego w instrukcji przed zakupem, bo nie jest to regułą.

Ta topologia ma dwie przewagi. Pierwsza dotyczy sterowania: skoro oba tory są niezależne, logika magazynu potrzebuje tylko informacji o bieżącym zużyciu i może sterować ładowaniem oraz rozładowaniem bez oglądania się na to, co robi mikrofalownik. W układzie szeregowym sytuacja jest znacznie trudniejsza – oprogramowanie musi rozdzielić strumień z mikrofalownika na część idącą do ładowania i część idącą do sieci, i autor odnotowuje, że część producentów sobie z tym nie radzi, a ich urządzenia wykazują zachowanie dwustanowe: albo ładują, albo oddają.

Druga przewaga jest mocowa. Przy naładowanym magazynie mikrofalownik i magazyn mogą zasilać instalację jednocześnie, dając do $2 \times 800 \text{ W}$. Scherer podaje przykład: przy odbiorniku 2 kW – grzałce pralki – z sieci trzeba dobrać tylko 400 W .

I tu pojawia się paradoks, który jest jednym z ciekawszych fragmentów oryginalnego artykułu: rozwiązanie technicznie najlepsze jest w Niemczech niedozwolone. Limit 800 W obowiązuje tam bez wyjątków i dotyczy sumy mocy elektrowni oraz magazynu, więc układ równoległy

kwalifikuje się jako zwykła instalacja fotowoltaiczna, ze wszystkimi tego konsekwencjami – projektem, uprawnionym wykonawcą, procedurą przyłączeniową. Autor argumentuje, że skoro do sieci publicznej i tak nigdy nie trafiłoby więcej niż 800 W , nie powinno mieć znaczenia, co dzieje się wewnątrz mieszkania, o ile zachowane są zasady zabezpieczenia obwodów – ale, jak sam kwituje, z przepisami się nie dyskutuje. Do polskiego statusu tej konfiguracji wrócimy w punkcie 11; jest inny.

3.4. Bilans strat

Przy porównywaniu topologii trzeba konsekwentnie rozdzielać dwa źródła strat, które w materiałach handlowych bywają mieszane: straty przetwarzania w przetwornicach i straty cyklu w samych ogniwach.

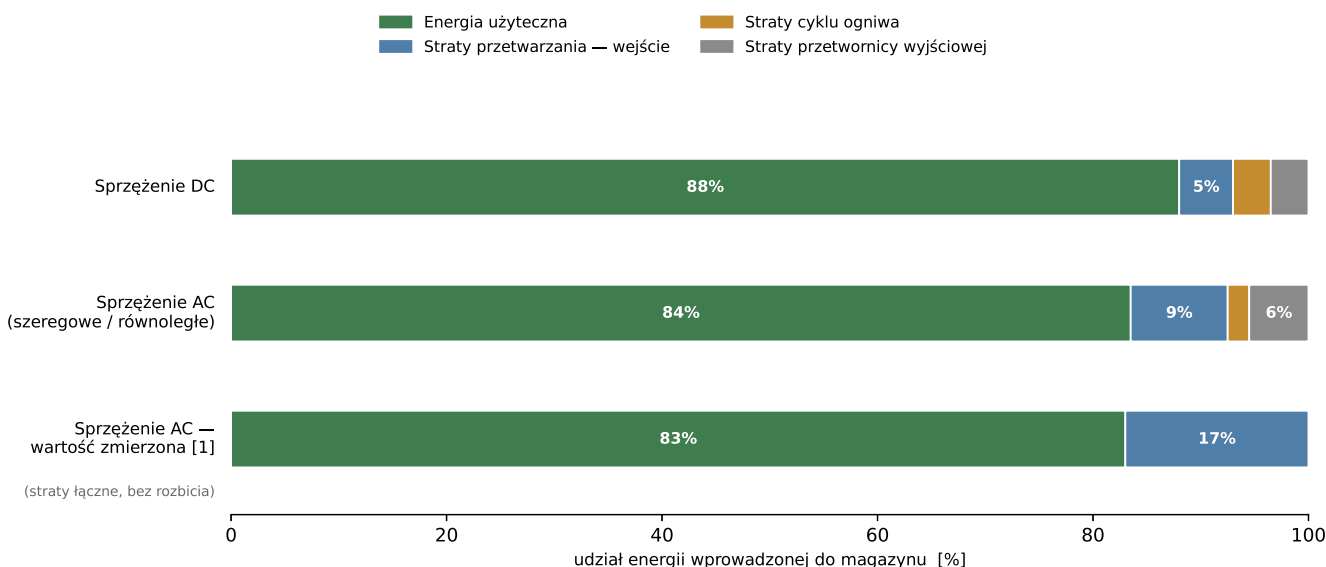
Sprawność cyklu (round-trip) stosowanych tu praktycznie bez wyjątku ogniów LiFePO_4 mieści się w przedziale $95\ldots98\%$, co daje $2\ldots5\%$ strat. To wartość, której nie zmienimy doбором topologii – jest własnością chemii ogniwa.

Straty przetwarzania są natomiast wprost funkcją liczby konwersji:

Sprzężenie DC. Jedna przetwornica na wyjściu, sprawność przetwarzania powyżej 90% , straty poniżej 10% . Po doliczeniu strat cyklu daje to sprawność całkowitą rzędu $85\ldots90\%$ i straty łączne $10\ldots15\%$.

Sprzężenie AC szeregowe. Dobry mikrofalownik osiąga $95\ldots96\%$, wbudowana przetwornica AC/DC podobnie – razem $8\ldots10\%$ strat jeszcze przed wejściem do ogniwa. Do tego straty cyklu (ponad 2%) i przetwornicy wyjściowej ($5\ldots6\%$), co daje straty łączne rzędu $15\ldots18\%$.

Scherer zaznacza, że są to wartości teoretyczne i że w praktyce urządzenia dostępne na rynku wypadają gorzej. Tym cenniejsza jest liczba zmierzona przez niego na własnym egzemplarzu magazynu ze sprzężeniem



Rysunek 5. Bilans strat dla obu topologii, z rozdzieleniem strat przetwarzania i strat cyklu ogniwa. Trzeci słupek to wartość zmierzona przez autora [1] na egzemplarzu ze sprzężeniem AC. Opracowanie własne, dane liczbowe wg [1]

Tabela 1. Porównanie topologii magazynu przy elektrowni balkonowej 800 W

Kryterium	Sprzężenie DC	AC szeregowo	AC równoległe
Liczba konwersji przed ogniwami	1 (DC/DC, MPPT)	2 (DC/AC + AC/DC)	2 (DC/AC + AC/DC)
Straty przetwarzania	<10%	8...10% (wejście) + 5...6% (wyjście)	jak szeregowo
Straty całkowite (teoret.)	10...15%	15...18%	15...18%
Wymagany mikrofalownik zewn.	nie	tak	tak
Miejsce montażu magazynu	zewnątrz, przy modułach	wewnątrz	wewnątrz
Przewód od modułów/falownika	solarny 6 mm ²	zasilający 3×1,5 mm ²	zasilający 3×1,5 mm ²
Grzałka ogniw	zwykle konieczna	zbędna	zbędna
Trudność sterowania	średnia	wysoka	niska
Maks. moc do instalacji	800 W	800 W	2×800 W
Dopuszczalność w Niemczech	tak	tak	nie
Dopuszczalność w Polsce	tak, po zgłoszeniu	tak, po zgłoszeniu	tak, po zgłoszeniu (por. pkt 11)

AC: **ok. 17% strat całkowitych**, co przy powyższych założeniach jest wynikiem dobrym.

3.5. Kompromis temperaturowy

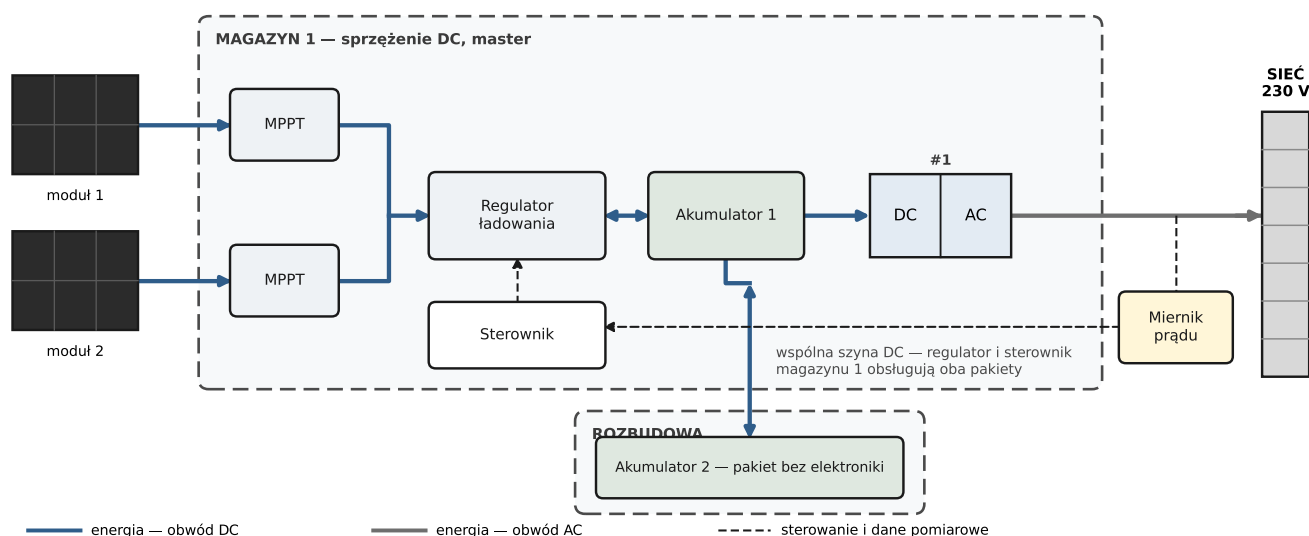
Zestawienie sprawności sugerowałoby przewagę sprzężenia DC. Praktyka jest bardziej złożona i decyduje o niej temperatura.

Ogniwa litowo-żelazowo-fosforanowe źle znoszą chłód – ładowanie przy temperaturze poniżej zera jest dla nich destrukcyjne, dlatego mniejsze magazyny przeznaczone do pracy na zewnątrz są niemal bez wyjątku wyposażane w grzałkę. Scherer wskazuje dwa problemy związane z tym rozwiązaniem. Po pierwsze, grzałka pobiera energię akurat zimą, czyli wtedy, gdy uzysk z modułów jest najmniejszy i każda kilowatogodzina jest zdobyta najwyższym kosztem – marnujemy więc na ogrzewanie baterii to, o co przez cały krótki dzień walczyliśmy. Po drugie, przy temperaturach wyraźnie poniżej zera grzałka i tak przestaje wystarczać.

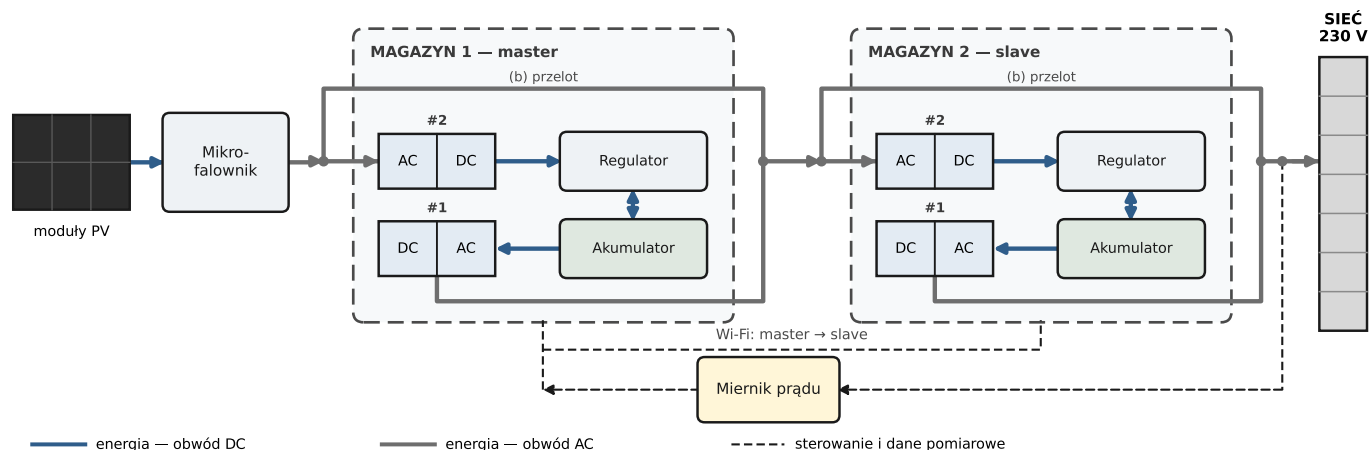
Sprzężenie AC pozwala tego uniknąć, bo magazyn stoi w ogrzewanym wnętrzu. Bilans wygląda więc tak: sprzężenie DC ma lepszą sprawność nominalną, ale zimą część tej przewagi oddaje na ogrzewanie własnej baterii; sprzężenie AC ma sprawność gorszą o kilka punktów procentowych, ale stabilną przez cały rok i wygodniejszy montaż. Dla instalacji w bloku wybór jest w praktyce przesądzony przez samą możliwość ustawienia urządzenia w mieszkaniu.

4. Rozbudowa pojemności

Pytanie o rozbudowę pojawia się naturalnie po pierwszym sezonie, zwłaszcza u użytkowników, którzy rozbudowali elektrownię z dwóch do czterech modułów. Ma to sens, choć nie zawsze jest oczywiste: cztery moduły o mocy 400...500 Wp każdy dają 1,6...2 kWp wejścia przy niezmiennym wyjściu 800 W. Na pierwszy rzut oka to przewymiarowanie, w rzeczywistości – sposób na utrzymanie pełnej mocy wyjściowej także w warunkach niekorzystnych: przy zachmurzeniu, złej orientacji



Rysunek 6. Rozbudowa pojemności przy sprzężeniu DC. Druga bateria dołączona po stronie stałoprądowej; sterowanie pozostaje w magazynie nr 1. Opracowanie własne na podstawie [1]



Rysunek 7. Rozbudowa przy sprzężeniu AC szeregowym. Magazyn nr 1 pracuje jako master, komunikuje się z miernikiem prądu i przekazuje polecenia magazynowi nr 2 przez Wi-Fi. Opracowanie własne na podstawie [1]

lub częściowym zaciemnieniu. Uzysk przy słabym świetle podwaja się, co zimą znaczy więcej niż latem.

Scherer podaje jako kryterium celowości rozbudowy magazynu średnie zapotrzebowanie powyżej 4 kWh na dobę przy czterech zainstalowanych modułach. Same warianty rozbudowy zależą od topologii.

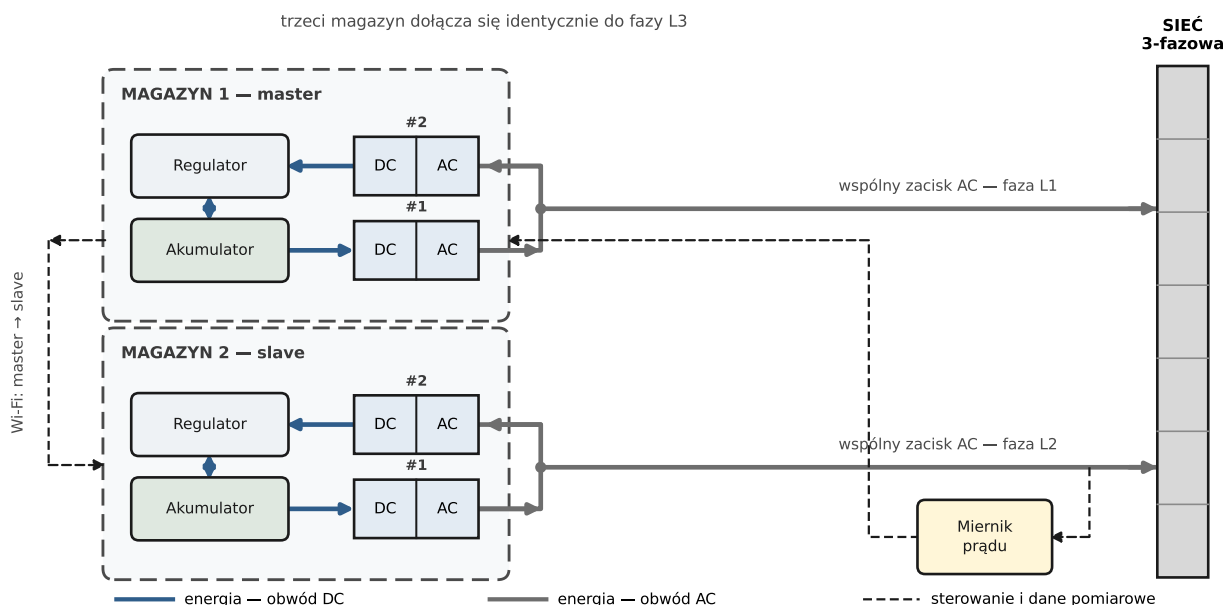
Przy sprzężeniu DC (rysunek 6) dokłada się drugi magazyn, którego elektronika jest w większości wyłączana, albo – jeżeli producent to przewiduje – samą baterię bez elektroniki. Magazyn nr 1 przejmuje pełne sterowanie i zarządza ogniwami obu urządzeń.

Przy sprzężeniu AC szeregowym (rysunek 7) dwa magazyny łączy się kaskadowo, przy czym pierwszy pracuje jako master, komunikuje się z miernikiem prądu i przekazuje polecenia drugiemu przez Wi-Fi. Producenci ograniczają zwykle taką kaskadę do dwóch urządzeń na jednej fazie, co przekłada się na praktyczne ograniczenie pojemności do 4...5 kWh.

Przy sprzężeniu AC równoległym (rysunek 8) kolejne magazyny można przyłączyć do różnych faz. To rozwiązuje przy okazji problem asymetrii obciążenia i podnosi moc dostępną dla instalacji: trzy magazyny to ponad 6 kWh pojemności i do 3,2 kW mocy. Autor sam uznaje to za przesadę jak na elektrownię balkonową, ale odnotowuje, że przy instalacji zarejestrowanej jako zwykła fotowoltaika i mocy modułów rzędu 8 kWp pojemność powyżej 6 kWh zaczyna mieć sens.

Warto przytoczyć też krytykę, którą Scherer formułuje pod adresem obecnych implementacji kaskady: bateria pierwszego magazynu jest ładowana do pełna i dopiero potem zaczyna się ładowanie drugiego. Autor zgadza się z użytkownikami, którzy uważają takie zachowanie za nieinteligentne – równomierne rozłożenie cyklu na oba pakiety byłoby korzystniejsze dla ich żywotności.

Skala kosztów przy rozbudowie: pojedyncze urządzenia kosztują 400...600 €, więc zestaw sześciu magazynów



Rysunek 8. Rozbudowa przy sprzężeniu AC równoległym. Kolejne magazyny przyłączone do osobnych faz, co jednocześnie ogranicza asymetrię obciążenia. Opracowanie własne na podstawie [1]

to ok. 3000 €, czyli ok. 250 €/kWh. Jak na magazyn stacjonarny nie jest to wynik zły – dla porównania pełnowymiarowe magazyny domowe wysokonapięciowe plasują się w przedziale kilkukrotnie wyższym.

5. Doświadczenia eksploatacyjne

Wartość artykułu Scherera polega w dużej mierze na tym, że opisuje konkretny egzemplarz w konkretnym budynku, a nie kartę katalogową.

Autor wybrał magazyn ze sprzężeniem AC – Hoymiles HiBattery 1920 AC o pojemności 1,92 kWh i masie 24 kg, w cenie 530 €. Wcześniejszy model tego samego producenta, MS-A2, ma pojemność 2,24 kWh, masę 35 kg i kosztował 580 €; różnica 11 kg przy porównywalnej pojemności jest tu argumentem niebagatelnym, bo urządzenie trzeba powiesić na ścianie. O wyborze przesądziła też perspektywa rozbudowy identycznymi urządzeniami – starszy model prawdopodobnie zniknie z oferty.

Montaż jest przykładem rozsądnego minimalizmu: magazyn zawisł w kotłowni, ok. 40 cm pod stropem, w miejscu, w którym do budynku wchodzi przewód od mikrofalownika. Zamiast firmowego uchwyty autor użył perforowanych blach stalowych z marketu budowlanego i nie musiał wiercić nowych otworów. Urządzenie ma dwa przyłącza – dolne od mikrofalownika, górne do sieci – a producent dopuszcza szeregowie połączenie najwyżej dwóch takich magazynów.

Z obserwacji ruchowych najciekawsze jest zachowanie regulatora. Wewnętrzne sterowanie koryguje moc wyjściową co 1...2 sekundy, dzięki czemu moc na przyłączy sieciowym utrzymuje się zwykle w zakresie pojedynczych watów. Scherer wprost pisze, że jest to wynik lepszy, niż się spodziewał – i jest to, dodajmy od siebie, najlepszy dowód na to, że zerowe wprowadzanie przestało być hasłem marketingowym.

Z liczb eksploatacyjnych warto zapamiętać dzień 5 lutego 2026, kiedy magazyn po raz pierwszy naładował się do pełna. Przy dopuszczalnym rozładowaniu do 10% stanu naładowania oznaczało to uratowanie 1,728 kWh energii przed oddaniem do sieci, co przy cenie 0,32 €/kWh daje 0,55 €. Sam autor uznaje tę wartość za zaniżoną, bo tego dnia dom nie był pusty, a ekspres, czajnik i zmywarka pobierały krótkotrwale ok. 2 kW. Magazyn działał więc dodatkowo jako bufor szczytów – piki poboru z sieci spłaszczyły się do ok. 1,2 kW – i według historii z aplikacji producenta rzeczywista oszczędność wyniosła tego dnia ok. 2,3 kWh, czyli 0,73 €.

Prognoza roczna, którą autor przedstawia z wyraźnym zastrzeżeniem, że jest szacunkiem, a nie pomiarem, opiera się na obciążeniu bazowym gospodarstwa: latem ok. 150 W, średniorocznie ok. 250 W. Dla krótkich nocy letnich daje to ok. 8 h × 180 dni × 200 W ≈ 288 kWh, dla

półroczna zimowego szacunkowo czwartą część tej wartości, czyli ok. 72 kWh. Do tego dochodzi efekt buforowania większych obciążeń dziennych, latem istotny, zimą pomijalny.

I tu trzeba postawić granicę między artykułem oryginalnym a tym, co z niego wynika dla polskiego czytelnika. Wszystkie powyższe kwoty są przeliczane po niemieckiej cenie detalicznej, ponieważ u Scherera każda kilowatogodzina zatrzymana w domu zastępuje kilowatogodzinę kupioną. W polskim net-billingu to założenie nie obowiązuje i cały rachunek trzeba zbudować od nowa. Robimy to w części drugiej tego artykułu.

Część II. Polska specyfika

6. Dlaczego niemiecki rachunek u nas nie działa

Cała ekonomiczna motywacja artykułu Scherera opiera się na jednym założeniu: nadwyżka wyprodukowanej energii, jeśli nie zostanie zużyta na miejscu, jest bezpowrotnie tracona. Przy elektrowni balkonowej w Niemczech nie zawiera się umowy o rozliczaniu energii wprowadzanej do sieci, a licznik z blokadą wsteczną nie zapisuje wprowadzenia. Każda kilowatogodzina zatrzymana w akumulatorze zastępuje zatem kilowatogodzinę kupioną po cenie detalicznej – i to jest cena, po której autor liczy oszczędności.

W Polsce mechanizm jest inny. Prosument objęty net-billingiem oddaje nadwyżkę do sieci, jej wartość trafia na depozyt prosumencki wyceniany według rynkowej ceny energii (RCE), a od nowelizacji z listopada 2024 r. wartość depozytu jest dodatkowo powiększana o współczynnik korekcyjny 1,23. Nadwyżka nie przepada – jest tylko wyceniana istotnie taniej niż energia kupowana, bo cena detaliczna zawiera dystrybucję, opłaty i podatki, których w RCE nie ma.

Konsekwencja dla opłacalności magazynu jest zasadnicza. Zysk z zatrzymania kilowatogodziny w domu to u nas nie pełna cena detaliczna, lecz **różnica** między nią a wyceną nadwyżki. Ta różnica jest oczywiście dodatnia i nadal znacząca, ale sprawia, że magazyn ma do odrobienia znacznie mniej, niż wynikałoby z niemieckiego rachunku – a jego cena jest ta sama.

Nie jest to spekulacja: rachunek został policzony dla polskich warunków. Szymon Sendlak i Jan Sitnik z Instytutu Sterowania i Elektroniki Przemysłowej Politechniki Warszawskiej opublikowali we wrześniowym numerze Przeglądu Elektrotechnicznego z 2025 r. analizę techniczno-ekonomiczną instalacji balkonowej dla warunków warszawskich [2]. Założenia są bliskie temu, o czym piszemy: dwa moduły po 400 W zamontowane pionowo na balustradzie skierowanej na południe, bez zacienienia, produkcja 643 kWh DC rocznie i 578 kWh po stronie AC przy sprawności przetwarzania 90%, zużycie

Tabela 2. Wyniki ekonomiczne dla instalacji balkonowej 800 W w Warszawie wg [2] (horyzont 10 lat, stopa dyskontowa 3%, wzrost cen energii 5% rocznie)

Scenariusz	Nakłady	SPBT [lata]	NPV 10 lat	IRR 10 lat
S2 – PV bez magazynu, bez oddawania do sieci	3000 zł	5,92	3 395 zł	16%
S3 – PV bez magazynu, net-billing	2500 zł	4,22	4 947 zł	26%
S4 – PV + magazyn 5 kWh, bez oddawania do sieci	9000 zł	11,24	1 229 zł	3%
S5 – PV + magazyn 5 kWh, taryfa dynamiczna z arbitrażem	8500 zł	7,21	6 423 zł	12%

gospodarstwa dwuosobowego 2031 kWh rocznie, opcjonalny magazyn 5 kWh o mocy 2,5 kW. Produkcję wyznaczono w PV*SOL, profil zużycia – w LoadProfileGenerator. Autorzy rozpatrzyli pięć scenariuszy, od układu odniesienia bez instalacji po instalację z magazynem współpracującym z rynkiem.

Wniosek autorów jest jednoznaczny i dla czytelnika przyzwyczajonego do niemieckiej narracji zaskakujący: **najkrótszy okres zwrotu daje instalacja pracująca w net-billingu bez magazynu**. Rozwiązanie to jednocześnie maksymalizuje wykorzystanie wyprodukowanej energii i minimalizuje nakłady początkowe. Dołożenie magazynu, którego koszt stanowi ok. 70% nakładów całego wariantu, wydłuża zwrot do ponad jedenastu lat i sprowadza wewnętrzną stopę zwrotu do 3%, a więc poniżej progu sensowności. Autorzy dodają istotne ostrzeżenie: przy ograniczonej żywotności ogniw istnieje realne ryzyko degradacji magazynu przed odzyskaniem zainwestowanych środków.

Nie oznacza to, że Scherer się myli – oznacza, że jego rozumowanie jest poprawne w jego systemie prawnym i błędne w naszym. Gdyby przenieść jego instalację do Polski bez zmiany sposobu sterowania, magazyn straciłby większość swojego uzasadnienia ekonomicznego. Pytanie brzmi więc: czy istnieje polskie uzasadnienie? Istnieje, ale zupełnie inne.

7. Kiedy magazyn się broni

Odpowiedź kryje się w scenariuszu S5 tej samej pracy. Jest to układ z magazynem, który nie tylko buforuje energię słoneczną, ale również dynamicznie współpracuje z rynkiem: ładuje się, gdy hurtowa cena energii brutto spada poniżej progu, i oddaje energię do sieci, gdy przekracza próg górny. Autorzy wyznaczyli te progi optymalizacyjnie, minimalizując roczny koszt energii gospodarstwa, i otrzymali odpowiednio 0,50 zł/kWh dla ładowania i 1,61 zł/kWh dla oddawania.

Wynik: NPV 6423 zł i IRR 12% w horyzoncie dziesięcioletnim – najlepszy rezultat spośród wszystkich wariantów z magazynem i wyższa wartość bieżąca netto niż w czystym net-billingu bez magazynu, choć przy dłuższym okresie zwrotu. Analiza wrażliwości pokazuje przy tym, że jest to wynik stabilny: w każdym wariantcie

zmiany nakładów, stopy dyskontowej i tempa wzrostu cen najlepiej wypadają scenariusze S3 albo S5.

Istotne zastrzeżenie autorów: **samo przejście na taryfę dynamiczną nie generuje znaczących oszczędności**. Korzyść pojawia się dopiero wtedy, gdy profil zużycia gospodarstwa zostanie dostosowany – zużycie przesunięte na godziny wysokiej generacji z PV lub niskich cen rynkowych, ograniczone w sytuacji przeciwniej. Magazyn jest tu narzędziem, które takie przesunięcie realizuje automatycznie, bez udziału mieszkańców. Autorzy odnotowują też, że opłacalność magazynów może rosnąć w miarę zwiększania się zmienności cen na rynku – trend obserwowany w pierwszej połowie 2025 r., z wyraźnie większą liczbą godzin z cenami ujemnymi.

Wniosek dla polskiego użytkownika można sformułować krótko. Ładowanie nadwyżkowe i zerowe wprowadzanie, czyli funkcje, które w Niemczech decydują o sensie zakupu, u nas są przydatne, ale nie rozstrzygające. **Rozstrzygające jest to, czy magazyn potrafi ładować się z sieci w godzinach tanich i oddawać w drogich** – a więc czy obsługuje taryfy dynamiczne i pozwala zaprogramować progi cenowe. To zupełnie inne kryterium zakupowe niż w niemieckich testach i warto o tym pamiętać, czytając zachodnie recenzje. Co ciekawe, jest to dokładnie ta funkcja, którą Scherer opisuje jako marginalną i odkłada na bok jako niezwiązaną z tematem swojego artykułu.

Warto też odnotować konsekwencję konstrukcyjną: w topologii ze sprzężeniem DC ładowanie z sieci wymaga dodatkowej przetwornicy AC/DC na wejściu, w topologiach AC – wewnętrzna przetwornica ma bezpośrednią ścieżkę do sieci i funkcja jest dostępna praktycznie zawsze. Dla polskiego użytkownika jest to więc kolejny, tym razem ekonomiczny argument za sprzężeniem AC.

8. Formalności – czego naprawdę wymagają polskie przepisy

Tu trzeba zacząć od obalenia twierdzenia, które w polskim internecie powtarzane jest z uporem: że instalacja do 800 W jest z czegokolwiek zwolniona. Polskie prawo nie zna kategorii „elektrownia balkonowa” i nie zna niemieckiego progu bagatelnego.

Status prawny instalacji. Każde źródło fotowoltaiczne przyłączone do instalacji budynku, która jest

połączona z siecią, jest mikroinstalacją w rozumieniu ustawy o OZE [11] – próg wynosi 50 kW i 800 W mieści się w nim tak samo jak 8 kW. Przyłączenie następuje na podstawie zgłoszenia do operatora systemu dystrybucyjnego, składanego co do zasady na 30 dni przed planowanym przyłączeniem. Procedura jest bezpłatna, a wymiana licznika na dwukierunkowy leży po stronie operatora. Zgłoszenie jest jednocześnie warunkiem uzyskania statusu prosumenta, a więc i rozliczenia nadwyżek – bez niego instalacja nie tylko funkcjonuje nieprawidłowo formalnie, ale też traci to, co w polskich warunkach stanowi o jej opłacalności.

Jedyny przypadek, w którym zgłoszenie nie jest potrzebne, to instalacja rzeczywiście wyspowa, fizycznie niezdolna do wprowadzenia energii do instalacji połączonej z siecią. Zestaw wpięty do gniazdka takim układem nie jest.

Zgoda właściciela nieruchomości. Balustrada, ściana zewnętrzna i elewacja stanowią część wspólną nieruchomości w rozumieniu ustawy o własności lokali [14], niezależnie od tego, że sama powierzchnia balkonu służy wyłącznie właścicielowi lokalu. Zamocowanie modułu do barierki albo przeprowadzenie przewodu przez elewację jest zatem ingerencją w część wspólną i wymaga zgody zarządu wspólnoty lub – w spółdzielni – jej organu. W praktyce zarządcy oczekują dodatkowo dokumentacji: sposobu mocowania, kart katalogowych, oświadczenia elektryka, coraz częściej polisy OC. To najczęstszy punkt, na którym instalacje w blokach się zatrzymują.

Prawo budowlane. Urządzenia fotowoltaiczne o mocy do 150 kW są zwolnione z obowiązku uzyskania pozwolenia na budowę na podstawie art. 29 ust. 2 pkt 16 [13]. Próg 6,5 kW, od którego wymagane jest uzgodnienie projektu z rzeczoznawcą do spraw zabezpieczeń przeciwpożarowych, przy zestawie balkonowym nie ma

zastosowania. Nie zwalnia to jednak z obowiązku wykonania mocowania zgodnie ze sztuką budowlaną – moduł zawieszony na wysokości nad terenem ogólnodostępnym jest elementem, którego oderwanie stwarza zagrożenie, a obciążenia wiatrowe na elewacji budynku wysokiego bywają istotnie większe niż na dachu domu jednorodzinnego.

Wymagania dla mikrofalownika. Urządzenie musi mieć certyfikat zgodności z wymaganiami NC RfG (rozporządzenie Komisji UE 2016/631 [16]) – jest to dokument, którego operator wymaga przy zgłoszeniu – oraz działające zabezpieczenie przed pracą wyspową. Na rynku wciąż występują tańsze zestawy z falownikami bez takiego certyfikatu; ich przyłączenie do sieci 230 V jest niedopuszczalne.

Magazyn energii. Domowy magazyn nie wymaga pozwolenia na budowę. Podlega natomiast zgłoszeniu do operatora – przy mikroinstalacji do 50 kW jest to zgłoszenie aktualizacji już istniejącej mikroinstalacji. Wymagane są dane magazynu i falownika, schemat połączeń oraz oświadczenie instalatora. Uwagi wymagają ograniczenia relacji mocy i pojemności magazynu względem mikroinstalacji wynikające z art. 7 ust. 8d12 Prawa energetycznego [12]; w praktyce bywają one powodem odrzucania zgłoszeń i są przedmiotem sporu w branży.

Pułapka statusu prosumenta. Status prosumenta dotyczy energii wytworzonej w odnawialnym źródle. Magazyn ładowany z sieci w godzinach tanich, a następnie oddający energię do sieci, wykracza poza ten schemat – energia oddawana nie pochodzi wówczas z OZE. Praktycznym rozwiązaniem jest właśnie tryb zerowego wprowadzania: magazyn oddaje wyłącznie tyle, ile zużywa instalacja budynku, więc do sieci nie trafia nic i problem nie powstaje. Warto to zestawić z punktem 7: arbitraż cenowy w polskich warunkach powinien być realizowany po stronie unikniętego zakupu, a nie sprzedaży do sieci. Magazyn ładuje się

Formalności krok po kroku

Kolejność czynności przy zestawie balkonowym 800 W z magazynem energii:

- 1. Zgoda wspólnoty lub spółdzielni** na zamocowanie modułów do balustrady lub elewacji. Przygotować sposób mocowania, karty katalogowe, oświadczenie elektryka; wiele wspólnot oczekuje też polisy OC.
- 2. Dobór osprzętu z certyfikatem NC RfG** – dotyczy mikrofalownika. Tańsze zestawy bez certyfikatu nie mogą być przyłączone do sieci 230 V.
- 3. Wykonanie instalacji.** Zalecany wydzielony obwód i przyłączenie na stałe zamiast wtyczki; przy konfiguracji równoległej – obowiązkowo (por. rozdz. 9).
- 4. Zgłoszenie mikroinstalacji wraz z magazynem do OSD**, co do zasady 30 dni przed przyłączeniem. Procedura bezpłatna. Wymagane: dane modułów, falownika i magazynu, schemat połączeń, certyfikat NC RfG, oświadczenie instalatora.

- 5. Wymiana licznika na dwukierunkowy** – po stronie operatora. Od tego momentu obowiązuje rozliczenie w net-billingu.

Uwaga. Mimo powszechnych twierdzeń w internetowych poradnikach nie znaleźliśmy przepisu, który zwalniałby instalację o mocy do 800 W z obowiązku zgłoszenia. Polskie prawo nie zna progu bagatelnego analogicznego do niemieckiego. Zwolniona jest wyłącznie instalacja wyspowa, niezdolna do wprowadzenia energii do sieci.



tanio i zasila dom w godzinach drogich – to jest wariant bezpieczny formalnie i to on odpowiada za większość korzyści w scenariuszu S5.

9. Wtyczka do gniazdka – najstabszy punkt konstrukcji

To jest zagadnienie, którego w polskim piśmiennictwie praktycznie nie ma, a które z punktu widzenia elektryka jest najważniejsze w całym temacie.

Instalacja gniazd wtykowych w mieszkaniu jest projektowana jako obwód promieniowy zasilany z jednej strony – z rozdzielnicy, przez zabezpieczenie nadprądowe, zwykle 16 A. Zabezpieczenie to chroni przewód obwodu przed skutkami przetężenia, przy czym cała logika doboru opiera się na założeniu, że prąd wpływa do obwodu wyłącznie przez to zabezpieczenie.

Wpięcie źródła w dowolny punkt takiego obwodu to założenie łamie. Powstaje obwód zasilany dwustronnie, w którym odcinek przewodu między punktem wstrzykiwania prądu a odbiornikiem może być obciążony sumą prądu z rozdzielnicy i prądu ze źródła – a wyłącznik nadprądowy w rozdzielnicy tej sumy nie widzi, bo mierzy tylko swoją część. Przy jednym mikrofalowniku 800 W margines jest wystarczający: dodatkowe ok. 3,5 A na obwodzie 16 A mieści się w zapasie wynikającym z obciążalności długotrwałej przewodu. Sytuacja zmienia się jednak przy rozbudowie. W konfiguracji równoległej z punktu 3.3, przy jednoczesnej pracy mikrofalownika i magazynu, mówimy o 2×800 W, czyli ok. 7 A wstrzykiwanych w obwód, a przy trzech magazynach na trzech fazach – o 3,2 kW.

Że nie jest to problem wyłącznie teoretyczny, sygnalizuje praktyka niemiecka: według doniesień prasowych sądy krajowe zakazały sprzedaży magazynów łączonych z systemami balkonowymi o mocy przekraczającej 960 W, jeżeli nie mają one zabezpieczeń chroniących przewody

instalacji przed przeciążeniem. Jest to o tyle znamienne, że dotyczy kraju, w którym przyłączanie elektrowni balkonowej przez gniazdo jest uregulowane i powszechnie akceptowane.

W Polsce takiego uregulowania nie ma w ogóle, co paradoksalnie stawia użytkownika w gorszej sytuacji: nie ma normy produktowej, na którą mógłby się powołać, ani wyraźnego zakazu. Rozwiązaniem zgodnym ze sztuką jest przyłączenie na stałe, do wydzielonego obwodu wyprowadzonego z rozdzielnicy mieszkania i wyposażonego we własne zabezpieczenie dobrane z uwzględnieniem pracy dwustronnej. Odsyłamy tu do PN-HD 60364-7-712 dotyczącej fotowoltaicznych układów zasilania [17] oraz PN-HD 60364-8-2 poświęconej niskonapięciowym instalacjom prosumenta [18] – ta druga jest w kontekście instalacji balkonowych zdecydowanie zbyt rzadko przywoływana.

Osobno warto zaznaczyć kwestię zabezpieczenia różnicowoprądowego. Przy źródle wprowadzającym prąd w obwód gniazd należy sprawdzić, czy zastosowany wyłącznik różnicowoprądowy jest właściwego typu dla prądów upływu o składowej stałej, które mogą pojawić się przy przetwornicach beztransformatorowych. Jest to standardowe zagadnienie przy instalacjach PV, ale przy zestawie „plug and play” kupowanym w markecie zwykle nikt go nie podnosi.

10. Wątek sieciowy – argument, o którym się nie mówi

Dotychczasowy rachunek był wykonany z perspektywy właściciela mieszkania. Istnieje jednak druga perspektywa, w której magazyn ma uzasadnienie zupełnie niezależne od rachunku za prąd, a której w popularnej dyskusji o fotowoltaice balkonowej praktycznie nie ma.

Zespół z Wydziału Elektrotechniki i Automatyki oraz Wydziału Fizyki Technicznej i Matematyki Stosowanej Politechniki Gdańskiej opublikował w styczniu 2026 r. w Energies symulację pracy sieci osiedla mieszkaniowego wykonaną w DiGSILENT PowerFactory [3]. Obiektem badań było rzeczywiste, nowo wybudowane osiedle 32 budynków wielorodzinnych w województwie pomorskim, po cztery lokale w budynku, o mocy przyłączeniowej 36 kW na budynek, zasilane z jednej stacji transformatorowej 15/0,4 kV czterema liniami kablowymi.

Wyniki dla wariantu, w którym każdy budynek wyposażono w fotowoltaikę o mocy równej jego mocy przyłączeniowej, przy minimalnym obciążeniu: przekroczenie dopuszczalnego poziomu napięcia na końcowych odcinkach linii między 9.30 a 14.30 oraz obciążenie transformatora sięgające **117%** w godzinach 10.15–14.15. W wariancie skrajnym, z mocą PV odpowiadającą sumie mocy przyłączeniowych lokali (50 kW na budynek), przeciążenie transformatora dochodzi do **176%**, a obciążenie linii

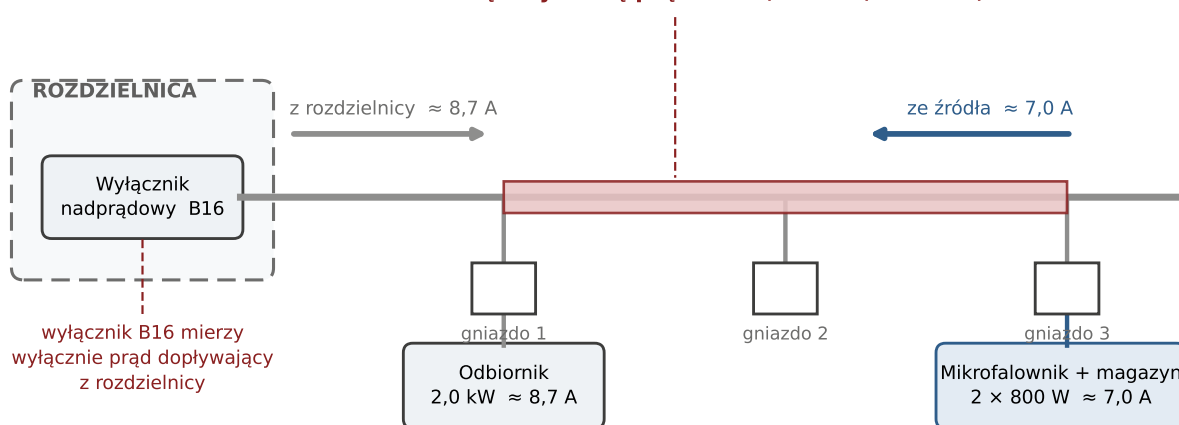
Ostrożnie ze źródłami internetowymi

Fotowoltaika balkonowa jest tematem, w którym pole informacyjne zostało zdominowane przez treści tworzone pod pozycjonowanie i linki afiliacyjne. Przygotowując ten artykuł, przejrzeliśmy kilkanaście polskojęzycznych „poradników 2026” i wnioski są jednoznaczne.

Serwisy te wzajemnie sobie przeczą w kwestii podstawowej – jedne twierdzą, że zestawu do 800 W nie trzeba zgłaszać do operatora, inne, że mimo uproszczonych przepisów zgłoszenie jest obowiązkowe. Część publikuje dane statystyczne przypisywane URE, których nie potwierdza żadne oficjalne źródło. Zdarzają się teksty publikowane po polsku w domenie zagranicznej, firmowane osobą doradcy energetycznego z doświadczeniem wyłącznie na rynku niemieckim i austriackim, przetykane odnośnikami z prowizją.

Wniosek praktyczny dla czytelnika: przy tym temacie warto sięgać do źródeł pierwotnych – ustawy, normy, karty katalogowej producenta, publikacji recenzowanej – albo do stanowiska własnego operatora sieci, uzyskanego na piśmie. To ostatnie jest w obecnym stanie prawnym najpewniejszą drogą, bo praktyka poszczególnych OSD bywa niejednorodna.

odcinek obciążony sumą prądów: $8,7 \text{ A} + 7,0 \text{ A} \approx 15,7 \text{ A}$



Zabezpieczenie w rozdzielni nie chroni odcinka przewodu między źródłem a odbiornikiem — prąd wstrzykiwany nie przepływa przez wyłącznik.

Rysunek 9. Obwód gniazd wtykowych zasilany dwustronnie. Odcinek między gniazdem ze źródłem a gniazdem z odbiornikiem przenosi sumę obu prądów, której zabezpieczenie w rozdzielni nie widzi. Wartości dla konfiguracji równoległej z rysunku 4

kablowych do 110%. Zmienność napięcia dobowego w węzłach sieci osiąga 0,16 j.w. przy teoretycznym maksimum 0,2 j.w. wynikającym z dopuszczalnego zakresu $\pm 10\%$.

Autorzy formułują wniosek, który wart jest zacytowania w streszczeniu: obecnie obowiązujące przepisy dotyczące budowy sieci zasilających osiedla wielorodzinne nie są skorelowane z przepisami dopuszczającymi montaż odnawialnych źródeł energii w tych budynkach. Sieć zaprojektowana zgodnie z normami dla samego poboru nie radzi sobie z sytuacją, którą te same przepisy dopuszczają po stronie wytwarzania.

Wskazują przy tym trzy możliwe kierunki działania. Pierwszy to odgórne ograniczenie mocy instalowanej – w badanym przypadku granicą było 32,5 kW na budynek, przy którym transformator pracuje w szczycie generacji dokładnie na 100% obciążenia. Drugi to przewymiarowanie elementów sieci, kosztowne i prowadzące do nadmiarowości względem rzeczywistych potrzeb. Trzeci – i to jest punkt istotny dla naszego tematu – **wdrożenie magazynów energii lub natychmiastowe lokalne zużycie nadwyżki**, na przykład przez konwersję na ciepło użytkowe.

Autorzy zwracają też uwagę na to, jak problem napięciowy jest rozwiązywany dzisiaj: przez zabezpieczenie nadnapięciowe falownika, które przy zbyt wysokim napięciu po stronie AC po prostu odłącza instalację. Jest to rozwiązanie działające, ale dla właściciela wybitnie niekorzystne, bo obniża uzysk dokładnie w godzinach najlepszej generacji.

Zastrzeżenie metodologiczne: praca dotyczy instalacji dachowych i fasadowych o mocach rzędu dziesiątków kilowatów na budynek, nie zestawów balkonowych 800 W. Mechanizm jest jednak identyczny, a różnica leży w rozproszeniu, nie w naturze zjawiska. Osiedle wielkopłytowe, w którym po zestawie balkonowym ma kilkuset lokatorów – a taki

potencjał opisuje analiza budynków wielkopłytowych opublikowana w *Energies* w listopadzie 2025 r., wskazująca, że w tego typu zabudowie mieszka ok. 30% populacji kraju [4] – wytworzy w południe sumaryczną moc porównywalną z instalacją dachową, tyle że rozłożoną przypadkowo na fazy i bez jakiegokolwiek koordynacji.

Wniosek jest następujący: magazyn przy elektrowni balkonowej ma uzasadnienie nie tylko w rachunku właściciela, ale też w bilansie lokalnej sieci niskiego napięcia. To argument, który dziś nie przekłada się jeszcze na żadną zachętę ani obowiązek, ale który – sądząc po kierunku wniosków cytowanej pracy – może się na nie przełożyć.

11. Konfiguracja równoległa: u nas nie jest zakazana

Wróćmy do paradoksu z rozdziału 3.3. Scherer opisuje układ równoległy jako technicznie najlepszy i jednocześnie zakazany, bo w Niemczech limit 800 W obowiązuje bezwzględnie i obejmuje łącznie elektrownię i magazyn.

W Polsce ten limit nie istnieje. Granicą jest moc przyłączeniowa lokalu oraz obowiązek zgłoszenia mikroinstalacji, a nie sztywna liczba 800 W. Oznacza to, że konfiguracja z rysunku 4, w Niemczech nielegalna, u nas jest dopuszczalna – pod warunkiem prawidłowego zgłoszenia zestawu wraz z magazynem i przy zachowaniu mocy poniżej mocy przyłączeniowej, co przy sumie rzędu 1,6 kW nie jest ograniczeniem istotnym.

Jest to różnica na korzyść polskiego użytkownika i warto ją wykorzystać świadomie, bo daje dostęp do topologii o najprostszym sterowaniu i najwyższej dostępnej mocy chwilowej. Z dwoma zastrzeżeniami. Pierwsze jest formalne: zgłoszenie musi obejmować rzeczywisty układ, a nie samą elektrownię. Drugie jest techniczne i dotyczy jednoczesnej pracy obu źródeł, problem obwodu zasilanego dwustronnie przestaje być teoretyczny.

Konfiguracja równoległa i przyłączenie przez gniazdo wychodzący się nawzajem; przy tym układzie wydzielony obwód przestaje być zaleceniem, a staje się warunkiem.

12. Podsumowanie i kryteria wyboru

Magazyn energii przy elektrowni balkonowej przeszedł w ciągu roku drogę od gadżetu do urządzenia, które da się obronić technicznie. Przełomem było sterowanie oparte na zewnętrznym pomiarze prądu – ładowanie nadwyżkowe i zerowe wprowadzanie – a nie postęp w ogniwach czy przetwornicach.

Dla polskiego użytkownika obraz wygląda jednak inaczej niż dla niemieckiego. Analiza z Politechniki Warszawskiej [2] pokazuje jednoznacznie, że sam magazyn buforujący energię słoneczną w polskim net-billingu wydłuża zwrot z ok. 4 do ponad 11 lat. Uzasadnieniem zakupu jest u nas

nie ratowanie nadwyżki przed przepadnięciem – bo ta nie przepada – lecz przesunięcie zużycia w czasie względem cen: ładowanie w godzinach tanich, zasilanie domu w drogich. To zmienia kryteria wyboru sprzętu.

Na co patrzeć przy zakupie:

- **Topologia pod kątem miejsca montażu.** W bloku sprzężenia AC jest wyborem praktycznie przesądzonym: magazyn stoi w mieszkaniu, nie potrzebuje grzałki i wymaga tylko przewodu zasilającego. Kilka punktów procentowych gorsza sprawność jest ceną, którą warto zapłacić.
- **Obsługa taryf dynamicznych i programowalne progi cenowe.** To jest w polskich warunkach funkcja rozstrzygająca o opłacalności. Warto sprawdzić, czy urządzenie pozwala zdefiniować progi liczbowo, a nie tylko wybrać tryb z listy.

Słowniczek akronimów i oznaczeń

Skróty i pojęcia użyte w artykule.

Akronimy i skróty

AC – prąd przemienny (*alternating current*); w tekście także obwód lub strona przemiennoprądowa układu

AC/DC, DC/AC – przetwornica prostująca (z przemiennego na stały) i przetwornica falująca, czyli falownik (ze stałego na przemienny)

BESS – baterijny magazyn energii (*battery energy storage system*)

CAN – magistrala komunikacyjna Controller Area Network, stosowana w większych instalacjach PV

DC – prąd stały (*direct current*); także obwód lub strona stałoprądowa układu

EMS – system zarządzania energią (*energy management system*)

IRR – wewnętrzna stopa zwrotu (*internal rate of return*) – stopa dyskontowa, przy której NPV inwestycji wynosi zero

LAN – lokalna sieć komputerowa, tu: przewodowa sieć domowa

LFP, LiFePO₄ – ogniwo litowo-żelazowo-fosforanowe – chemia stosowana praktycznie bez wyjątku w magazynach tej klasy

LOM – zabezpieczenie przed pracą wyspową (*loss of mains*) – odłącza źródło po zaniku napięcia sieci

MaStr – Marktstammdatenregister – niemiecki rejestr instalacji wytwórczych; polskim odpowiednikiem funkcjonalnym jest zgłoszenie do OSD

MPPT – układ nadążnego śledzenia punktu mocy maksymalnej modułu (*maximum power point tracking*)

NC RfG – kodeks sieci dotyczący wymogów przyłączania jednostek wytwórczych (Network Code Requirements for Generators), rozporządzenie Komisji (UE) 2016/631

NPV – wartość bieżąca netto (*net present value*) – zdyskontowana suma przepływów pieniężnych w przyjętym horyzoncie

OSD – operator systemu dystrybucyjnego

OZE – odnawialne źródła energii

PV – fotowoltaika, fotowoltaiczny (*photovoltaics*)

RCE – rynkowa cena energii elektrycznej – godzinowa cena stosowana do wyceny nadwyżek w net-billingu

RS485 – standard interfejsu transmisji szeregowej, stosowany w komunikacji falownik–licznik

SPBT – prosty okres zwrotu nakładów (*simple pay-back time*), bez dyskontowania

STC – standardowe warunki pomiarowe modułu PV (*standard test conditions*): 1000 W/m², 25 °C, AM 1,5

URE – Urząd Regulacji Energetyki

Wi-Fi – bezprzewodowa sieć lokalna

Pojęcia branżowe

ładowanie nadwyżkowe – (*surplus charging*) – tryb, w którym magazyn pobiera energię wyłącznie z nadwyżki produkcji ponad bieżące zapotrzebowanie budynku

zerowe wprowadzanie – (*zero feed-in*) – tryb, w którym magazyn oddaje do instalacji dokładnie tyle mocy, ile jest w danej chwili zużywane, bez oddawania nadmiaru do sieci

sprawność cyklu – (*round-trip efficiency*) – iloczyn sprawności ładowania i rozładowania magazynu, bez strat przetwarzania w przetwornicach

net-billing – obowiązujący w Polsce sposób rozliczenia prosumenta, w którym nadwyżka wprowadzona do sieci jest wyceniana wartościowo, a nie odbierana w naturze

depozyt prosumencki – konto, na którym gromadzona jest wartość energii wprowadzonej do sieci, przeznaczona na rozliczenie przyszłych zakupów

praca wyspowa – praca źródła na wydzieloną instalację po zaniku napięcia sieciowego; niedopuszczalna bez odpowiednich zabezpieczeń

mikroinstalacja – w rozumieniu ustawy o OZE instalacja odnawialnego źródła o mocy zainstalowanej nie większej niż 50 kW

Symbole i jednostki

W, kW – wat, kilowat – moc czynna

Wp, kWp – watopik, kilowatopik – moc szczytowa modułu PV w warunkach STC

Wh, kWh – watogodzina, kilowatogodzina – energia; kWh/rok w odniesieniu rocznym

V, A – volt, amper – napięcie i prąd

mm² – przekrój znamionowy żyły przewodu

B16 – wyłącznik nadprądowy o charakterystyce wyzwalania B i prądzie znamionowym 16 A

L1, L2, L3 – przewody fazowe instalacji trójfazowej

j.w. – jednostki względne (*per unit*) – wielkość odniesiona do wartości znamionowej

η – sprawność przetwarzania lub sprawność cyklu, wyrażana w procentach

- **Tryb zerowego wprowadzania i wymagany do niego miernik prądu.** Bez tego magazyn nie umie sterować się rozsądnie. Trzeba przy tym doliczyć koszt miernika i sprawdzić, czy magazyn go obsługuje.
- **Otwartość protokołu.** Lokalne API lub MQTT zamiast wyłącznej zależności od chmury producenta. To kryterium będzie zyskiwać na znaczeniu wraz z wiekiem urządzenia.
- **Certyfikat NC RfG mikrofalownika.** Warunek zgłoszenia i przyłączenia. Bez niego nie ma statusu prosumenta, a więc nie ma opłacalności.
- **Deklarowana liczba cykli.** Przy magazynie, który w polskich warunkach zwraca się w kilku–kilkunastu latach, żywotność ogni w jest parametrem ekonomicznym, nie tylko technicznym.

Kiedy magazynu nie kupować. Jeżeli instalacja pracuje w klasycznym net-billingu bez taryfy dynamicznej, a gospodarstwo nie ma wysokiego obciążenia wieczornego, magazyn będzie w obecnych cenach urządzeniem o wątpliwym zwrocie. Same panele zwracają się szybciej. Sytuacja zmieni się, jeśli spadną ceny magazynów, wzrośnie zmienność cen rynkowych albo – co po lekturze pracy z Politechniki Gdańskiej [3] nie jest scenariuszem nierealnym – pojawią się zachęty lub wymogi po stronie operatorów.

JT, WM



Bibliografia:

Źródło inspiracji

[1] T. Scherer, Battery Storage for Balcony Power Plants. Small Batteries, Big Impact, Elektor, May & June 2026, s. 106–113.

Publikacje – Polska

- [2] S. Sendak, J. Sitnik, Techno-Economic Assessment of balcony-mounted photovoltaics in urban areas/Analiza techniczno-ekonomiczna wykorzystania fotowoltaiki balkonowej w przestrzeni miejskiej, Przegląd Elektrotechniczny, R. 101, nr 9/2025, s. 208–214, DOI 10.15199/48.2025.09.30.
- [3] R. Kowalak, D. Kowalak, K. Seklecki, L. S. Litzbarski, Challenges in the Legal and Technical Integration of Photovoltaics in Multi-Family Buildings in the Polish Energy Grid, Energies 2026, 19(2), 474, DOI 10.3390/en19020474.
- [4] Balcony Photovoltaics in Large-Panel Prefabricated Buildings as a Contribution to the Urban Energy Transition, Energies 2025, 18(21), 5789, DOI 10.3390/en18215789.
- [5] T. Wiśniewski, M. Pawlak, Beyond Subsidies: Economic Performance of Optimized PV-BESS Configurations in Polish Residential Sector, Energies 2025, 18(24), 6615, DOI 10.3390/en18246615.
- [6] K. J. Cieślak, Profitability Analysis of a Prosumer Photovoltaic Installation in Light of Changing Electricity Billing Regulations in Poland, Energies 2024, 17(15), 3618, DOI 10.3390/en17153618.
- [7] K. Seklecki, M. Olesz, M. Adamowicz, M. Nowak, L. S. Litzbarski, K. Balcarek, J. Grochowski, A Comprehensive System for Protection of Photovoltaic Installations in Normal and Emergency Conditions, Energies 2025, 18(7), 1749, DOI 10.3390/en18071749.
- [8] L. S. Litzbarski, K. Seklecki, M. Adamowicz, J. Grochowski, PV installations and the safety of residential buildings, Inżynieria Bezpieczeństwa Obiektów Antropogenicznych 2023, nr 4, s. 1–10, DOI 10.37105/iboa.177.

Publikacje zagraniczne

- [9] U. Spaeth, A. Popp, H. Fechtner, A. Cichoń, B. Schmülling, Self-consumption Optimization of Balcony Solar Systems Using a Load-controlled Battery Storage, 10th International Conference on Power and Energy Systems Engineering (CPSE), Nagoya 2023, DOI 10.1109/CPSE59653.2023.10303065.
- [10] D. L. Gerber, A. Ginsberg-Klemmt, L. Stoler, J. Shackelford, A. Meier, Barriers to Balcony Solar and Plug-In Distributed Energy Resources in the United States, Energies 2025, 18(8), 2132, DOI 10.3390/en18082132.

Akty prawne i normy

- [11] Ustawa z dnia 20 lutego 2015 r. o odnawialnych źródłach energii, z późn. zm.
- [12] Ustawa z dnia 10 kwietnia 1997 r. – Prawo energetyczne, art. 7 ust. 8d–8d12.
- [13] Ustawa z dnia 7 lipca 1994 r. – Prawo budowlane, art. 29 ust. 2 pkt 16.
- [14] Ustawa z dnia 24 czerwca 1994 r. o własności lokali.
- [15] Ustawa z dnia 27 listopada 2024 r. o zmianie ustawy o odnawialnych źródłach energii oraz niektórych innych ustaw (współczynnik korekcyjny depozytu prosumenckiego).
- [16] Rozporządzenie Komisji (UE) 2016/631 z dnia 14 kwietnia 2016 r. ustanawiające kodeks sieci dotyczący wymogów w zakresie przyłączenia jednostek wytwórczych do sieci (NC RfG).
- [17] PN-HD 60364-7-712 Instalacje elektryczne niskiego napięcia. Część 7-712: Wymagania dotyczące specjalnych instalacji lub lokalizacji. Fotowoltaiczne (PV) układy zasilania.
- [18] PN-HD 60364-8-2 Instalacje elektryczne niskiego napięcia. Część 8-2: Niskonapięciowe instalacje elektryczne prosumenta.
- [19] PN-EN 62109-1 Bezpieczeństwo konwerterów mocy stosowanych w fotowoltaicznych systemach energetycznych. Część 1: Wymagania ogólne.
- [20] PN-EN IEC 61730-1 Ocena bezpieczeństwa modułu fotowoltaicznego (PV). Część 1: Wymagania dotyczące konstrukcji.

Materiały producentów

- [21] Hoymiles, karty katalogowe magazynów HiBattery 1920 AC oraz MS-A2.
- [22] Shelly, karta katalogowa miernika Pro 3EM.
- [23] GoodWe, karta katalogowa systemu ESA-Athena (falownik 0,8 kW, bateria LFP 1,92 kWh, wejście PV 2,4 kW, cztery MPPT).

AI w pracowni elektronika konstruktora (4)

Trzy zadania, z którymi AI radzi sobie na urządzeniu

Głos, obraz, sygnały – trzy fundamentalnie różne klasy problemów. Każda wymaga innego podejścia do danych, innej architektury modelu i innego profilu sprzętowego. Nie istnieje jeden „model AI do embedded”. Istnieją za to trzy dobrze scharakteryzowane dziedziny, w których inżynierowie elektronicy świadomi tych różnic wdrażają systemy działające niezawodnie w produkcji.

Prowadzimy przez każdą z tych dziedzin oddzielnie: opisujemy właściwości danych, typowe architektury modeli, ograniczenia sprzętowe i – co najważniejsze dla praktyka – udokumentowane przypadki wdrożeń z konkretnymi parametrami wydajnościowymi.

1

GŁOS
Interfejsy głosowe i rozpoznawanie mowy na urządzeniu

4.1 Dlaczego głos jest trudniejszy, niż się wydaje

Rozpoznawanie mowy wydaje się intuicyjnie prostym zadaniem – ludzie robią to bezwiednie od urodzenia. Dla algorytmu ML jest to jednak jeden z trudniejszych problemów percepcji, szczególnie w warunkach embedded. Sygnał mowy jest wielowymiarowy, wysoce zmienny i zależy od setek czynników: akcentu, płci, wieku mówiącego, odległości od mikrofonu, poziomu tła akustycznego i chwilowego stanu fizjologicznego użytkownika.

Do tego dochodzi wymiar prywatności. Interfejs głosowy, który dosłownie słucha całą dobę i wysyła nagrania do chmury, budzi uzasadnione obawy użytkowników – i coraz częściej także regulatorów. Przetwarzanie głosu lokalnie, bez jakiegokolwiek transmisji, usuwa ten problem strukturalnie, nie regulaminowo.

4.2 Architektura systemu głosowego on-device

System rozpoznawania głosu działający w całości na MCU składa się z trzech niezależnych warstw, które muszą współpracować ze sobą w czasie rzeczywistym.

Warstwa Audio Front End (AFE) jest często niedoceniana jako element „infrastrukturalny”, tymczasem to właśnie jej jakość decyduje o tym, z jakim sygnałem mierzy się model ML. Redukcja szumu tła i beamforming – ukierunkowanie czułości mikrofonu na źródło dźwięku – mogą być realizowane sprzętowo (dedykowany procesor audio) lub programowo na MCU z FPU. Wybór

Tabela 9. Trzy warstwy systemu głosowego on-device. Opracowanie własne na podst.: Renesas/EE World 2025; Lin et al. 2024		
Warstwa	Funkcja	Typowe zasoby
Audio Front End (AFE)	Przetwarzanie wstępne: redukcja szumu, beamforming, normalizacja poziomu	DSP dedykowany lub MCU z FPU; < 5% CPU
Ekstrakcja cech (Feature Extraction)	Konwersja sygnału na reprezentację cech: MFCC, Mel-spektrogram, filterbank	MCU + FFT; 5–15% CPU; kilka kB RAM
Model klasyfikacji/ KWS	Detekcja słów kluczowych lub pełne NLU (Natural Language Understanding)	10–100 kB Flash; 20–80 kB SRAM; < 10 ms latencja

Keyword Spotting – jak to naprawdę działa

Urządzenie działa w dwóch fazach energetycznych:

FAZA 1: ALWAYS-ON (mikrowaty)

- Mały, skrajnie energooszczędny detektor (np. DSP sprzętowy) nasłuchuje sygnału
- Model „wake word” jest prosty: rozróżnia ciszę od aktywności mowy
- Pobór mocy: rząd mikrowatów; bateria wystarcza na miesiące

FAZA 2: ACTIVE (miliwaty, wyzwolona)

- Po wykryciu aktywności budzi się główny MCU
- Pełny model KWS lub NLU analizuje nagrane słowo/polecenie
- Latencja inferencji: < 10 ms
- Przetworzone polecenie trafia do systemu; MCU wraca do stanu w fazie 1

Kluczowa zaleta: dane audio nigdy nie opuszczają urządzenia.

Źródło: Lin et al. 2024 (arXiv 2403.19076); Renesas/EE World 2025

CASE STUDY: System VUI na MCU bez połączenia sieciowego (Renesas VUI ecosystem)

Problem: sterowanie głosowe urządzenia w języku naturalnym bez dostępu do chmury.

Rozwiązanie:

- Platforma: mały MCU ogólnego przeznaczenia (rodzina RA) lub ASSP RISC-V
- Narzędzie: DSpotter Modeling Tool (Cyberon) – GUI do budowy modeli KWS/NLU
- Wbudowane modele: 44+ języków, eliminacja etapu zbierania danych treningowych
- Opcje rozszerzenia: Voice Anti-Spoofing, Speaker ID, Audio Front End

Wynik: lokalne rozpoznawanie komend głosowych bez transmisji do chmury, działanie offline, RODO-zgodna architektura przetwarzania dźwięku.

Źródło: Renesas Electronics/EE World, Shaping the Future with Embedded AI, 2025

mający bezpośrednie przełożenie na precyzję rozpoznawania w hałaśliwym środowisku.

Ekstrakcja cech MFCC (Mel-Frequency Cepstral Coefficients) to standard dla KWS: kompresuje sygnał audio do wektora kilkunastu–kilkudziesięciu współczynników opisujących barwę dźwięku w sposób zbliżony do percepcji ludzkiego ucha. Pozwala to zredukować wymiar wejścia modelu o rzędy wielkości względem surowej próbki PCM – co jest kluczowe dla MCU z ograniczoną pamięcią SRAM.

2

WIZJA

Computer vision od kontroli jakości do autonomicznej nawigacji

4.3 Computer vision na embedded: skala możliwości

Computer vision na urządzeniach embedded ewoluowała w ciągu ostatnich pięciu lat od klasyfikacji binarnej (jest obiekt/nie ma obiektu) do pełnej detekcji obiektów w czasie rzeczywistym, segmentacji semantycznej i wykrywania anomalii wizualnych. Każde z tych zadań stawia inne wymagania pamięci i mocy obliczeniowej, co bezpośrednio przekłada się na wybór platformy sprzętowej.

4.4 FOMO-AD: detekcja anomalii wizualnych bez próbek błędnych

Klasyczne podejścia do wizyjnej kontroli jakości miały fundamentalne ograniczenie: trzeba było zebrać próbki dla każdej możliwej klasy wady. Wątpliwa jakość spawu, pęknięcie powłoki, różne warianty uszkodzeń obudowy – każda z tych usterek wymagała dziesiątek lub setek przykładowych zdjęć. W praktyce przemysłowej jest to rzadko wykonalne: produkcja nie dostarcza równej ilości wszystkich rodzajów wad.

Architektura FOMO-AD (Faster Objects, More Objects – Anomaly Detection), zaprezentowana przez Edge Impulse w 2024 roku, odwróciła to założenie. Model jest trenowany wyłącznie na przykładach poprawnych (normalnych). Wszystko, co odbiega od wzorca normalnego, jest flagowane jako anomalia. Nie trzeba z góry wiedzieć, jak wygląda wada – wystarczy pokazać modelowi, jak wygląda produkt prawidłowy.

Kluczowe parametry wydajnościowe FOMO-AD: na procesorze Cortex-M7 @ 480 MHz architektura osiąga 30 klatek na sekundę przy zużyciu poniżej 200 kB RAM. Na procesorze Cortex-M4F @ 156 MHz to 5 klatek na sekundę. Ten sam model działa na pełnym GPU (Macbook) z przepustowością 1000 fps – co ilustruje skalowalność

Tabela 10. Zadania computer vision na embedded – wymagania i zastosowania. Opracowanie własne na podst.: Renesas/EE World 2025; Edge Impulse 2024; Huang et al. 2025

Zadanie CV	Trudność	Min. platforma	Przykład zastosowania
Klasyfikacja obrazów (prostych)	Niska	MCU Cortex-M4	Sortowanie elementów, kontrola obecności
Visual Wake Words (VWW)	Niska-średnia	MCU Cortex-M4/M7	Detekcja obecności osoby w kadrze
Detekcja obiektów (FOMO)	Średnia	MCU Cortex-M7/NPU	Linia produkcyjna, kontrola jakości
Detekcja anomalii (FOMO-AD)	Średnia-wysoka	MCU + <200 kB RAM	Wykrywanie defektów bez próbek błędnych
Segmentacja semantyczna	Wysoka	MPU (RZ/V, i.MX 8)	Autonomiczna nawigacja, AGV w magazynie
Rozpoznawanie twarzy/biometria	Wysoka	MPU + ISP + NPU	Kontrola dostępu, personalizacja HMI

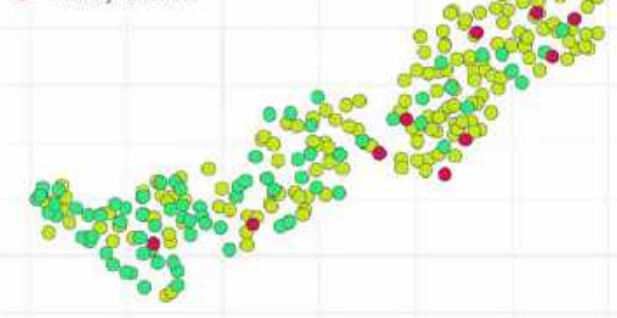
Model testing results

ACCURACY
95.90%

	ANOMALY	NO ANOMALY	UNCERTAIN
ANOMALY	94.1%	5.9%	-
NO ANOMALY	0%	100%	0%
F1 SCORE	0.97	0.94	

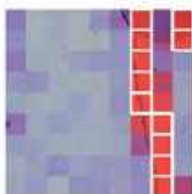
Feature explorer

● anomaly - correct
● no anomaly - correct
● anomaly - incorrect



anomaly.cracked...

Label: anomaly
Predicted: anomaly
[View sample](#)
[View classification](#)



Rysunek 4.1. FOMO-AD w akcji: detekcja pęknięcia w betonie (anomalii) przez model trenowany wyłącznie na danych normalnych.
Źródło: Edge Impulse 2024

architektury: identyczny model, radykalnie różne możliwości sprzętowe.

3

SYGNAŁY

Analityka czasu rzeczywistego – słuchanie maszyn i świata

4.5 Dane rokładów czasowych – dlaczego FFT nie wystarcza

Sygnały z czujników fizycznych – akcelerometrów, mikrofonów, mierników prądu/napięcia, sensorów temperatury – różnią się fundamentalnie od danych tabelarycznych czy obrazów. Charakteryzują się bardzo wysoką częstotliwością próbkowania, zjawiskiem jittera (nieregularna zmienność chwilowa sygnału) oraz sygnałami przejściowymi – krótkimi, niestacjonarnymi zdarzeniami, które zawierają kluczową informację o stanie maszyny.

Inżynierowie przez dekady pracowali z sygnałami za pomocą Szybkiej Transformaty Fouriera (FFT), która rozkłada sygnał na składowe częstotliwościowe. FFT jest cennym narzędziem – zwłaszcza dla urządzeń wirujących, gdzie piki przy wielokrotnościach częstotliwości obrotowej

FFT – zamieniamy sygnał w widmo częstotliwościowe, tracimy informację w dziedzinie czasu. Dane z akcelerometru uśrednione przez sekundę mówią nam głównie o grawitacji.

Renesas/Reality AI, Shaping the Future with Embedded AI, EE World 2025

są diagnostycznie interpretowalne. Jednak FFT ma jedną zasadniczą słabość: pokazuje całe widmo częstotliwościowe jako jeden wynik uśredniony w czasie, tracąc informację o chwilowych zdarzeniach.

Badania eksperymentalne z 2025 roku (PMC Sensors) dostarczają konkretnych danych o tej różnicy. Przy detekcji anomalii na danych wibracyjnych, akustycznych i środowiskowych porównano cztery podejścia do ekstrakcji cech: wyłącznie FFT, wyłącznie Wavelet, kombinacja obu oraz model bez transformacji. Wynik jest jednoznaczny: cechy oparte na transformacie falkowej (Wavelet) konsekwentnie przewyższały czyste FFT w metryce F1. Skwantyzowany model TinyML (INT8) osiągnął F1 $\approx 0,92$ przy 3-krotnej redukcji śladu pamięciowego i 60% niższym zużyciu energii względem pełnoprecyzyjnego odpowiednika.

4.6 Predictive maintenance: od częstotliwości do kondycji maszyny

Predyktywne utrzymanie ruchu (predictive maintenance) to zastosowanie, które najwyraźniej ilustruje wartość ekonomiczną Edge AI. Zamiast reagować na awarię (maintenance reaktywny) lub wymieniać części według stałego harmonogramu (maintenance planowy), system wykrywa narastającą degradację – i alarmuje z odpowiednim wyprzedzeniem, kiedy interwencja jest jeszcze możliwa bez przestoju produkcji.

Czujnik wibracji połączony z MCU i modelem TinyML może zostać zamontowany bezpośrednio na silniku lub łożysku. System działa autonomicznie, bez połączenia z siecią: zbiera dane wibracyjne, wydziela cechy z surowego sygnału i porównuje aktualny profil z zapisanym wzorcem „normalnej pracy”. Odchylenie powyżej progu generuje

CASE STUDY: Detekcja defektów wizualnych w linii produkcyjnej (FOMO-AD, 2024)

Problem: identyfikacja anomalii na powierzchni betonów, płytek, elementów metalowych bez możliwości zebrania kolekcji zdjęć każdego rodzaju wady.

Architektura FOMO-AD:

- Trening: wyłącznie próbki „normalne” (brak wad)
- Detekcja: Gaussian Mixture Model na mapach cech (GMM lub PatchCore)
- Wyjście: wynik anomalii per-region z lokalizacją na obrazie
- Platformy: MCU (Cortex-M7: 30 fps, <200 kB RAM) – GPU (1000 fps)

Zastosowania produkcyjne: kontrola ICB, spawy, powłoki, inspekcja baterii EV, kontrola pojemników farmaceutycznych.

Źródło: Edge Impulse Inc., FOMO-AD release 2024; BusinessWire 23.04.2024

Tabela 11. Porównanie metod detekcji anomalii na danych sensorycznych (wibracje, akustyka, środowisko). Źródło: PMC Sensors 2025, 25(21), 6629. doi: 10.3390/s25216629

Metoda ekstrakcji cech	Wynik F1	Pamięć skala względna	Energia skala względna
Shallow Neural Network (SNN) – pełna precyzja	0,94	100%	100%
TinyML (Quantized INT8) – FFT + Wavelet	0,92	33% (3× redukcja)	40% (60% oszczędności)
SNN – tylko FFT (bez Wavelet)	0,89	95%	95%
Decision Tree – FFT + statystyki	0,85	15%	20%

CASE STUDY: Detekcja łuków elektrycznych w instalacji PV/ESS – model na MCU

Problem: łuki elektryczne w panelach fotowoltaicznych/systemach ESS to zagrożenie pożarowe. Konwencjonalne algorytmy osiągają tylko ok. 50% dokładności detekcji.

Rozwiązanie (MCU RA6M4 + Reality AI Tools):

- Dane: sygnał prądowy z instalacji PV – sygnatury szumowe łuków
- Model: kompaktowy, automatycznie wygenerowany przez Reality AI Tools
- Rozmiar modelu: ≤ 8 kB RAM, 23 kB Flash
- Dokładność detekcji: ≥ 99,8%
- Skalowalność: obsługa prądów do 200 A bez utraty skuteczności

Znaczenie: ponad dwukrotna poprawa dokładności względem dotychczasowych metod, model działający lokalnie, bez dostępu do sieci.

Źródło: Renesas Electronics/EE World, Shaping the Future with Embedded AI, 2025.

alarm – lokalny, bez potrzeby transmisji. Badania Heydari i Mahmoud dokumentują implementację na pompie odśrodkowej, która osiągnęła dokładność detekcji anomalii na poziomie 99,9%.

4.7 Dlaczego trzy filary wymagają różnych decyzji projektowych

Połączenie trójki głos-wizja-sygnały w jednym rozdziale nie jest przypadkowe: razem tworzą przestrzeń możliwości multimodalnych urządzeń. Nowoczesne systemy embedded AI coraz częściej angażują więcej niż jedno źródło danych jednocześnie: rozpoznanie głosu aktywuje kamerę, która potwierdza tożsamość mówiącego; czujnik wibracji wykrywa anomalię, którą kamera następnie lokalizuje wizualnie.

Jednak każdy z trzech filarów rządzi się inną logiką decyzji projektowych. Głos wymaga niskiego poboru mocy i always-on listening; wizja – przepustowości pasma i pamięci na konwolucyjny model sieci neuronowej;

sygnały – próbkowania z wysoką częstotliwością i starannej ekstrakcji cech zachowującej informację z dziedziny czasu. Inżynier, który rozumie te różnice, wybiera właściwy sprzęt i właściwy algorytm zanim rozpocznie zbieranie danych – a nie po miesiącach eksperymentów.

Podsumowanie

Trzy filary AI na urządzeniu – głos, wizja, sygnały – to nie warianty tego samego problemu, lecz odrębne dziedziny z własną logiką. Głos wymaga always-on przy mikrowatach i przetwarzania chroniącego prywatność; architektura FOMO-AD pozwala na detekcję wad wizualnych trenując na samych przykładach normalnych, działając z 30 fps na Cortex-M7 w mniej niż 200 kB RAM; analiza sygnałów wymaga ekstrakcji cech zachowującej informację z dziedziny czasu (Wavelet > FFT), co umożliwia TinyML modele o F1 ≈ 0,92 przy 3-krotnej redukcji pamięci.

Wdrożenie produkcyjne – detekcja łuków elektrycznych w instalacji PV z dokładnością ≥ 99,8%, model w ≤ 8 kB RAM – pokazuje, że nie chodzi o akademickie demonstracje. Edge AI w tych trzech dziedzinach działa w produkcji, generując wymierną wartość ekonomiczną.

WM, JS

ŹRÓDŁA:

- [1] Renesas Electronics/EE World. Shaping the Future with Embedded AI. Digital Ebook 2025 – VUI (rozdz. AI/ML Voice), wizja DRP-AI3, case study AFCI/PV (rozdz. 5).
- [2] Edge Impulse Inc. FOMO-AD – Faster Objects, More Objects Anomaly Detection. Premiera kwiecień 2024; BusinessWire 23.04.2024; dokumentacja techniczna 2024.
- [3] Heydari S., Mahmoud Q.H. Tiny Machine Learning and On-Device Inference: A Survey. Sensors 2025, 25(10), 3191 – detekcja anomalii na pompie odśrodkowej (99,9%), aplikacje healthcare.
- [4] Lin J. et al. Tiny Machine Learning: Progress and Futures. IEEE CAS Magazine 2024. arXiv 2403.19076 – KWS, always-on, mapa zastosowań.
- [5] PMC Sensors 2025, 25(21), 6629. Lightweight Signal Processing and Edge AI for Real-Time Anomaly Detection in IoT Sensor Networks. doi: 10.3390/s25216629 – FFT vs Wavelet, TinyML F1=0,92, 3x redukcja pamięci.
- [6] Springer IoT 2025. IoT device for detecting abnormal vibrations in motors using TinyML. doi: 10.1007/s43926-025-00142-4 – wykrywanie usterek łożysk silnika, 96,5% dokładność, 300 ms latencja.

Tabela 12. Porównanie trzech filarów AI na urządzeniu – właściwości danych i priorytety sprzętowe. Opracowanie własne

Filar	Kluczowa właściwość danych	Główne wyzwanie	Priorytet sprzętowy
GŁOS	Wysoka zmienność mówcy; szum tła	Always-on przy mikrowatach; prywatność	Niski pobór mocy; DSP audio
WIZJA	Wysoka rozdzielczość; redundancja pikseli	Rozmiar modelu CNN; przepustowość	NPU/ISP; duża pamięć Flash
SYGNAŁY	Sygnały przejściowe; jitter; szum	Ekstrakcja cech; utrata info. z dziedziny czasu	FPU/DSP; w.cz. próbkowania ADC



Kod z autopilota

Programowanie mikrokontrolerów wspomagane przez AI – stan na jesień 2026

Firma badawcza RunSafe Security przepytala w grudniu 2025 roku ponad dwustu inżynierów oprogramowania wbudowanego z USA, Wielkiej Brytanii i Niemiec, pracujących nad urządzeniami dla infrastruktury krytycznej – a więc nad sterownikami sieci energetycznych, aparaturą medyczną, elektroniką samochodową i systemami sterowania w zakładach przemysłowych. Osiemdziesiąt trzy i pół procent z nich odpowiedziało, że w wyrobach, które ich firmy już sprzedały i które pracują dziś u klientów, znajduje się firmware zawierający fragmenty napisane przez model językowy. Nie prototypy, nie projekty wewnętrzne – wyroby na rynku.

W tym samym czasie firma Veracode zakończyła czwartą rundę pomiaru, w którym ponad sto modeli językowych dostaje zadania programistyczne dające się rozwiązać zarówno w sposób bezpieczny, jak i podatny na cyberatak, a wynik sprawdza analizator bezpieczeństwa kodu. Modele generują kod, który się kompiluje w niemal stu procentach przypadków. Kod wolny od podatności z listy OWASP Top 10 generują w pięćdziesięciu sześciu procentach przypadków – czyli dokładnie tak samo dobrze jak rok wcześniej i dwa lata wcześniej. Składnię modele opanowały. Bezpieczeństwa nie i od dwóch lat to się nie poprawia.

Te dwie liczby nie są ze sobą sprzeczne, choć zestawione obok siebie brzmią alarmująco. Opisują ten sam stan z dwóch stron: narzędzie weszło do warsztatu szybciej, niż warsztat

zdążył się do niego przystosować. To nie jest artykuł o tym, czy używać sztucznej inteligencji przy pisaniu firmware'u – ta decyzja zapadła już bez udziału większości z nas, i to również w firmach, które formalnie jej zakazały. To jest artykuł o tym, jak sprawdzić kod pochodzący z modelu, zanim trafi on do pamięci mikrokontrolera: jakie narzędzie wykrywa który rodzaj błędu, jaki kontekst trzeba modelowi dostarczyć, żeby przestał zgadywać, i za co – od 11 września 2026 roku – odpowiada producent wyrobu.

1. Punkt wyjścia: co się właściwie zmieniło

Przez cztery lata narzędzie zmieniło klasę trzy razy – i nie stało się tak dlatego, że modele zrobiły się mądrzejsze.



Rysunek 1. Trzy pokolenia narzędzi AI w programowaniu – i moment, w którym zmienia się rola inżyniera

Pierwsze pokolenie, z lat 2022–2023, to autouzupełnianie w edytorze. Model widział kilkaset linii wokół kursora i proponował następną. Odpowiedzialność za każdą linię pozostawała po stronie człowieka, bo człowiek każdą linię widział.

Drugie pokolenie, z lat 2024...2025, to czat świadomy projektu. Model dostawał kontekst całego repozytorium, potrafił wyjaśnić cudzy kod, zaproponować przebudowę, napisać test. Człowiek nadal był operatorem: kopiował, wklejał, uruchamiał kompilator i oceniał wynik.

Trzecie pokolenie, które ustabilizowało się w 2026 roku, to agent z dostępem do łańcucha narzędzi. Sam wywołuje polecenie budowania, sam wgrywa obraz na płytke, sam czyta bufor diagnostyczny RTT, sam interpretuje ostrzeżenia kompilatora i sam poprawia to, co znalazł. Różnica nie polega na jakości generowanego kodu, bo ta zmienia się powoli. Polega na tym, że człowiek przestał widzieć każdą linię, zanim ta trafi do repozytorium.

To rozróżnienie wraca w niniejszym artykule wiele razy, warto więc postawić je na początku wprost: **większość problemów opisanych dalej w tym artykule istniała już w 2023 roku. Zmieniła się jedynie liczba miejsc, na które nikt nie patrzy.**

Skala zjawiska w oprogramowaniu wbudowanym

Dane dotyczące naszej branży są bardziej konkretne, niż się zwykle sądzi, ponieważ powstały co najmniej dwa niezależne badania ankietowe skupione wyłącznie na systemach wbudowanych.

Wspomniane już badanie RunSafe Security objęło ponad dwustu specjalistów pracujących nad urządzeniami dla infrastruktury krytycznej. Osiedziesiąt koma pięć procent respondentów używa sztucznej inteligencji przy

tworzeniu oprogramowania. Osiedziesiąt trzy i pół procent ma kod pochodzący z modelu już w produkcji, przy czym blisko połowa deklaruje, że taki kod działa w więcej niż jednym z ich wyrobów. Żaden respondent nie zadeklarował, że unika tych narzędzi całkowicie.

Badanie przeprowadzone przez Censuswide na zlecenie firmy Black Duck objęło 785 specjalistów od oprogramowania wbudowanego i pokazało coś, co powinno zainteresować każdego, kto odpowiada w firmie za zgodność wyrobów z przepisami. Osiedziesiąt dziewięć procent organizacji używa asystentów AI. Trzydzieści dziewięć procent dopuszcza je tylko dla wybranych pracowników. Dwadzieścia pięć procent formalnie ich zakazuje – i w trzech czwartych spośród tych firm programiści używają ich mimo zakazu. Zjawisko doczekało się nazwy: *shadow AI*, czyli sztuczna inteligencja w cieniu (w szarej strefie).

Z perspektywy czysto inżynierskiej jest to ciekawostka socjologiczna. Z perspektywy producenta wprowadzającego wyrób na rynek Unii Europejskiej jest to kod nieznanego pochodzenia w produkcji, pod którym firma podpisuje deklarację zgodności. Wracamy do tego wątku w rozdziale poświęconym rozporządzeniu CRA.

Dlaczego oprogramowanie wbudowane jest inne

Przez pierwsze lata fali AI przyrost wydajności koncentrował się niemal wyłącznie w oprogramowaniu webowym, biurowym i mobilnym. Firmware został z tyłu i nie był to przypadek.

W aplikacji serwerowej istnieje przeogromny materiał do trenowania AI, a wąskie gardło rozwiązuje się do-daniem zasobów. W projektowaniu softwaru dla mikro-kontrolera zasoby materiałów do treningu AI są nader

Trzy warstwy użycia AI w oprogramowaniu wbudowanym

Warstwa 1 – asystent w edytorze. Podpowiada linie, uzupełnia komentarze, generuje szkielety funkcji. Człowiek widzi każdą propozycję przed jej przyjęciem. Ryzyko: zmyślane nazwy rejestrów i funkcji, kod stylistycznie poprawny i merytorycznie fałszywy. Minimum kontrolne: kompilacja z maksymalnym poziomem ostrzeżeń oraz analiza statyczna działająca w edytorze.

Warstwa 2 – agent w repozytorium. Czyta całą bazę kodu, wprowadza zmiany w wielu plikach naraz, dopisuje testy, dociąga zależności. Ryzyko: powielanie zamiast ponownego użycia, komponenty obce, o których nikt świadomie nie zdecydował, zmiany rozjeżdżające architekturę. Minimum kontrolne: przegląd nastawiony na architekturę i powielenia, kontrola zależności w potoku ciągłej integracji, generowanie wykazu komponentów przy każdym budowaniu.

Warstwa 3 – agent na stanowisku sprzętowym. Buduje, wgrywa, uruchamia, czyta logi i kanał RTT, poprawia na podstawie zachowania układu. Ryzyko: skasowanie pamięci urządzenia produkcyjnego, trwałe zablokowanie interfejsu debugowego przez wgrany program rozruchowy, iterowanie „aż zadziała” bez zrozumienia przyczyny. Minimum kontrolne: fizyczne oddzielenie stanowiska doświadczalnego od sprzętu produkcyjnego, jawna lista operacji nieodwracalnych wymagających potwierdzenia przez człowieka, pomiar wykonywany niezależnie od agenta.

skromne, a poza tym nie ma instancji, którą można doskonalować, ani sposobu na dołożenie pamięci, gdy jej zabraknie. Poprawność jest tu wielowymiarowa: kod może być funkcjonalnie bezbłędny i jednocześnie nie do przyjęcia, ponieważ przekracza budżet czasowy procedury obsługi przerwania, pobiera dwa miliampery ponad specyfikację albo nie mieści się w sekcji przewidzianej przez skrypt konsolidatora.

Do tego dochodzą rejestry sprzętowe, których rzeczywiste zachowanie bywa udokumentowane wyłącznie w erracie – wykazie błędów samego układu scalonego, publikowanym osobno pod nazwą errata sheet, oraz horyzont czasowy: firmware wgrany do wyrobu przemysłowego pracuje kilkanaście lat, a wycofanie zmiany oznacza kłopotliwą i kosztowną kampanię serwisową, a nie przywrócenie poprzedniej wersji w trybie wdrożeniowym.

Model językowy nie ma dostępu do żadnego z tych wymiarów, o ile mu się go celowo nie dostarczy. I to jest właściwa diagnoza – nie „sztuczna inteligencja nie radzi sobie z oprogramowaniem wbudowanym”, tylko „model pracuje

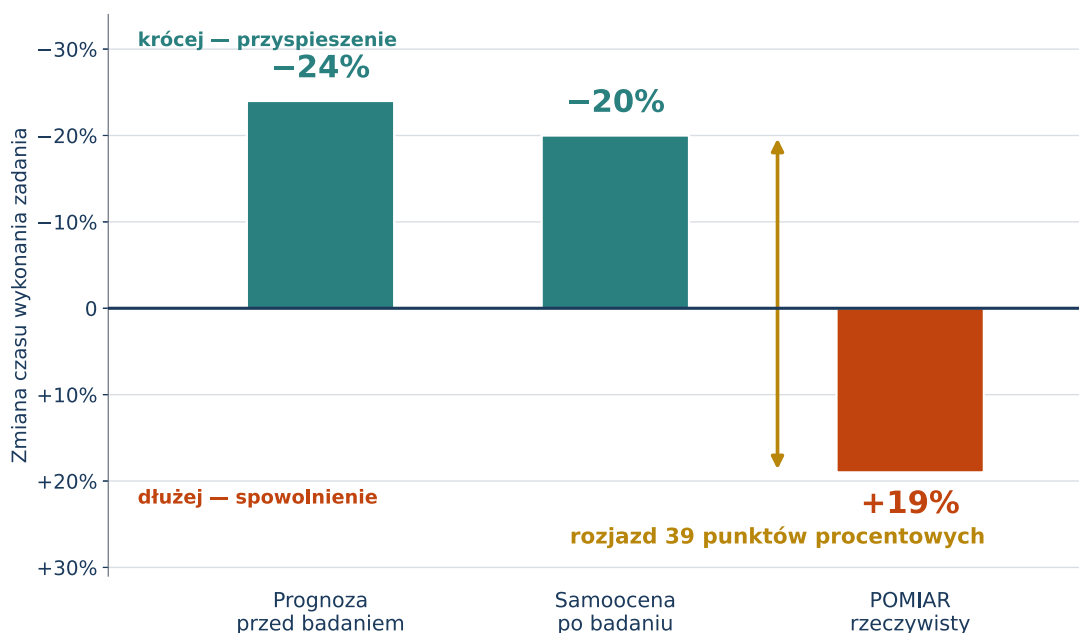
na tej części opisu systemu, którą mu pokazano, a w firmwarze pokazuje mu się zwykle najmniej istotny fragment”.

2. Ile z tego jest realnym zyskiem?

Zanim przejdziemy do techniki, potrzebna jest rzetelna podstawa liczbowa. W tej dyskusji anegdota wypiera pomiar wyjątkowo skutecznie, a materiały dostawców podają wskaźniki przyspieszenia rzędu trzydziestu do pięćdziesięciu procent bez metodologii, do której dałoby się zajrzeć.

Eksperyment, który wyszedł odwrotnie

W lipcu 2025 roku organizacja badawcza METR opublikowała wynik eksperymentu przeprowadzonego z randomizacją i grupą kontrolną, a więc metodą, którą stosuje się w badaniach klinicznych, nie zaś w ankietach branżowych. Szesnastu doświadczonych programistów pracujących nad projektami otwartymi wykonało 246 realnych zadań w dojrzałych repozytoriach, nad którymi pracowali średnio od pięciu lat. Każde zadanie



Rysunek 2. Badanie METR: 16 doświadczonych deweloperów, 246 realnych zadań. Uczestnicy byli przekonani, że przyspieszyli – pomiar wykazał spowolnienie

losowo przydzielano do wariantu „narzędzia AI dozwolone” albo „narzędzia AI zakazane”.

Przed rozpoczęciem uczestnicy prognozowali, że narzędzia skrócą czas pracy o dwadzieścia cztery procenty. Po zakończeniu oceniali, że skróciły go o dwadzieścia procent. Pomiar wykazał, że w rzeczywistości zadania wykonywane z pomocą AI trwały o dziewiętnaście procent dłużej.

Rozjazd między odczuciem a wynikiem wyniósł trzydzieści dziewięć punktów procentowych. Ludzie, którzy właśnie osobiście doświadczyli spowolnienia, byli przekonani, że przyspieszyli.

Powtórka, która się nie udała – i dlaczego to ważne

W lutym 2026 roku METR opublikował rzecz w badaniach rzadką: komunikat o tym, że dane z drugiej edycji własnego eksperymentu nie nadają się do wyciągania wniosków. Druga edycja objęła 57 programistów, 143 repozytoria i ponad osiemset zadań. Wyniki surowe sugerowały tym razem przyspieszenie – o osiemnaście procent w podgrupie uczestników pierwszego badania i o cztery procent wśród nowo zrekrutowanych, w obu przypadkach z przedziałami ufności przechodzącymi przez zero, a więc bez definitywnego rozstrzygnięcia co do kierunku.

Zespół uznał te liczby za niewiarygodne z powodu efektu doboru próby. Od trzydziestu do pięćdziesięciu procent uczestników przyznało, że przestali zgłaszać do badania zadania, których nie chcieli wykonywać bez wsparcia narzędzi. Jeden z uczestników opisał to wprost: unikam zgłaszania problemów, które AI zamknie w dwie godziny, a ja bez niej będę rozwiązywał dwadzieścia godzin. Przestała działać także rekrutacja – rosnąca grupa programistów odmawia udziału w badaniu, w którym

połowę pracy trzeba wykonać bez narzędzi, do jakich już przywykła.

Płynie stąd wniosek metodologiczny: **im bardziej przydatne jest narzędzie, tym trudniej zmierzyć jego wpływ**, ponieważ ludzie przestają zgadzać się na pracę bez niego. Każdy wskaźnik przyspieszenia publikowany od tej pory należy czytać z tym zastrzeżeniem, niezależnie od tego, w którą stronę wskazuje. Postawa programistów, którzy nie mogą już pracować bez użycia AI jest bardziej wymowna, niż jakiegokolwiek badania statystyczne.

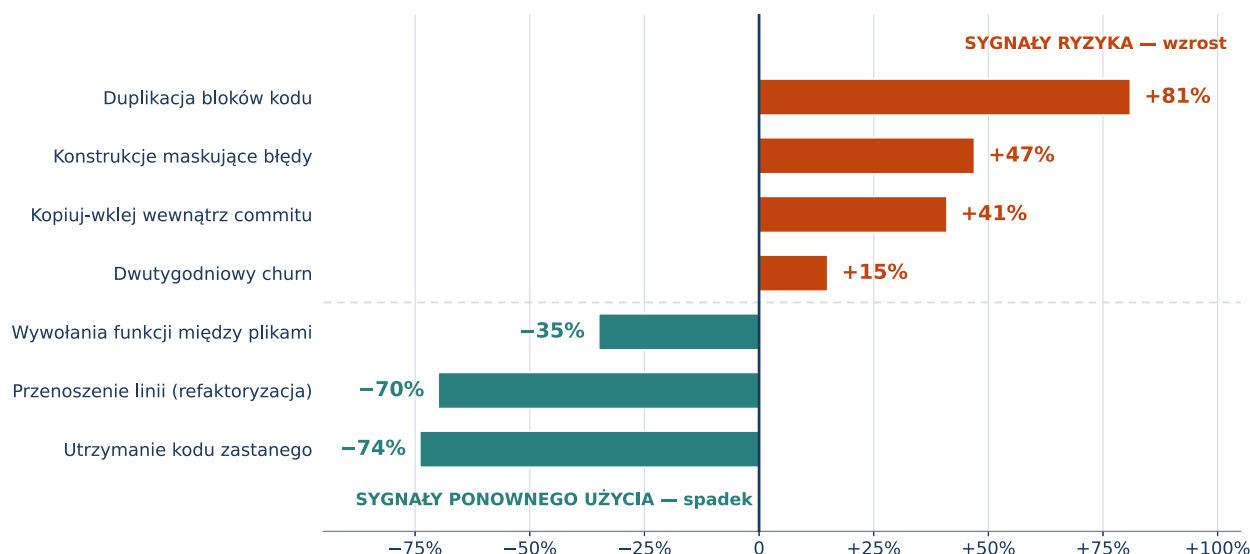
Przepustowość rośnie, utrzymywalność spada

Drugie źródło twardych danych ma zupełnie inną naturę – jest to analiza rzeczywistych zmian w kodzie, nie zaś deklaracji. Firma GitClear przeanalizowała 623 miliony zmian z lat 2023...2026, klasyfikując każdą zmienioną linię i śledząc osiem wskaźników jakości.

Wskaźniki ponownego użycia kodu spadły. Wywołania funkcji między plikami, świadczące o tym, czy nowy kod korzysta z tego, co już istnieje, zmalały o trzydzieści pięć procent. Przenoszenie linii przy przebudowie kodu – o siedemdziesiąt procent. Utrzymanie kodu zastanego – o siedemdziesiąt cztery procent wobec poziomu z 2022 roku.

Wskaźniki ryzyka wzrosły. Kopiowanie i wklejanie wewnątrz pojedynczej paczki zmian (commitu) – o czterdzieści jeden procent. Powielanie całych bloków kodu – o osiemdziesiąt jeden procent. Dwutygodniowy współczynnik przepisywania, czyli odsetek kodu usuwanego lub zmienianego w ciągu dwóch tygodni od wprowadzenia – o piętnaście procent.

Dla firmware’u najgroźniejszy jest jednak czwarty wskaźnik: konstrukcje maskujące błędy wzrosły o czterdzieści



Rysunek 3. Sygnały jakości kodu 2023...2026 wobec poziomu bazowego. Analiza 623 mln zmian: przepustowość rośnie, utrzymywalność spada

siedem procent. Pod tą nazwą kryją się: pusty blok obsługi wyjątku, zignorowany kod powrotu z funkcji biblioteki sprzętowej, pętla oczekiwania na flagę bez limitu czasu. W aplikacji serwerowej jest to rozwiązanie nieoptymalne. W sterowniku, który ma zareagować na zanik zasilania, jest to różnica między zadziałaniem a zawieszeniem.

Warto przy tym zauważyć, że ta sama firma opublikowała w styczniu 2026 roku wynik przemawiający w drugą stronę: użytkownicy intensywni wytwarzają od czterech do dziesięciu razy więcej kodu niż osoby nieużywające tych narzędzi. Po odjęciu efektu doboru – bo po nowe narzędzia sięgają najpierw ci najbardziej wydajni – realny przyrost wobec ich własnych wcześniejszych wyników wynosi około dwudziestu pięciu procent. To nadal dużo, ale nie jest to dziesięciokrotność.

Kto ma rację

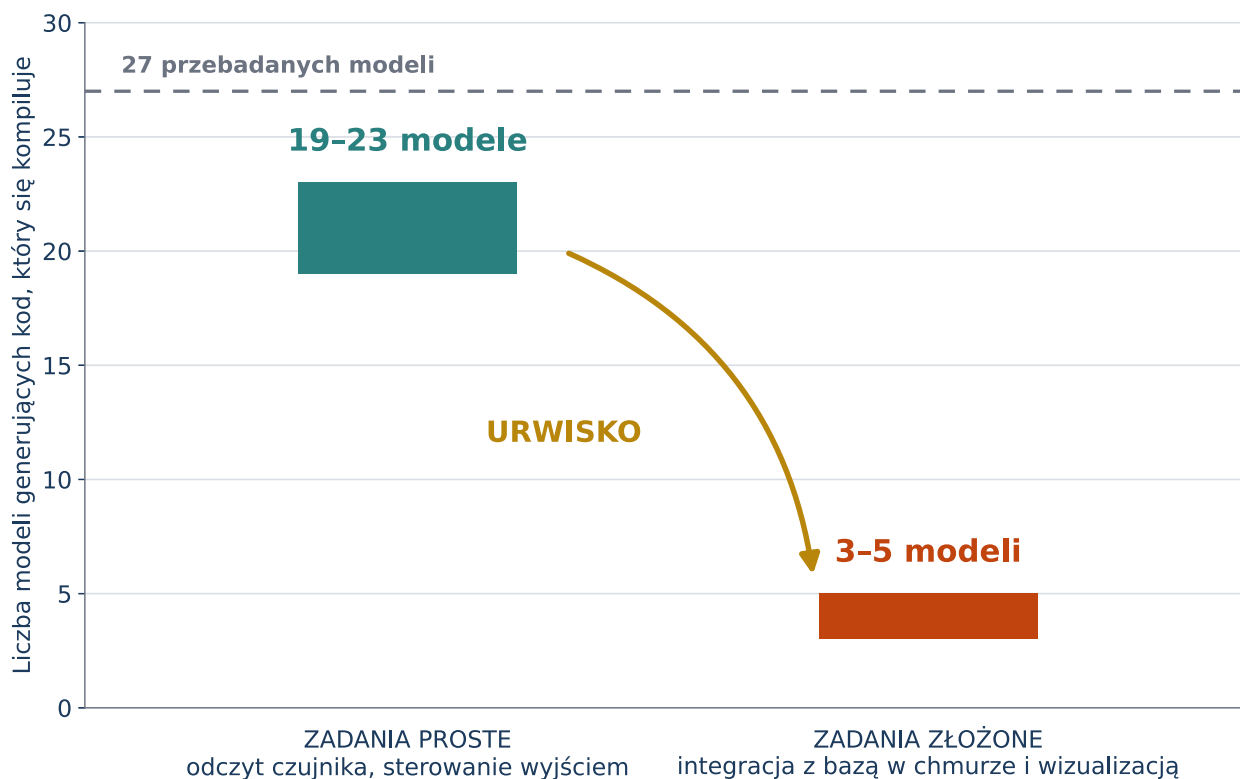
Sprzeczność między tymi źródłami jest pozorna i najlepiej rozbraja ją raport DORA, oparty na danych z tysięcy zespołów. Jego teza brzmi: sztuczna inteligencja jest wzmocniaczem. Wzmacnia mocne strony organizacji, która ma uporządkowany proces, i wzmacnia dysfunkcje organizacji, która go nie ma. W wydaniu z 2026 roku, poświęconym zwrotowi z inwestycji, dochodzi obserwacja jeszcze bardziej praktyczna: zysk pojawia się dopiero wtedy, gdy przegląd kodu i wewnętrzna infrastruktura nadążają za przyrostem wolumenu. Jeżeli nie nadążają, wąskie gardło nie znika – przesuwa się o jeden etap dalej i wraca jako zaległości w przeglądzie, poprawki i awarie na produkcji.

Potwierdza to badanie serwisu Stack Overflow z 2025 roku, w którym adopcja i zaufanie rozjechały się po raz pierwszy w historii tego pomiaru. Osiemdziesiąt cztery procent respondentów używa narzędzi AI lub planuje ich użycie, ponad połowa zawodowych programistów codziennie – a odsetek deklarujących zaufanie do trafności wyników spadł z czterdziestu do dwudziestu dziewięciu procent. Dwie trzecie badanych jako główny problem wskazuje odpowiedzi „prawie dobre, ale nie do końca”. Jest to najdroższa kategoria błędu, jaka istnieje: zbyt poprawna, by odrzucić ją na pierwszy rzut oka, i zbyt błędna, by program zadziałał prawidłowo.

Wniosek dla konstruktora jest następujący. Zysk czasowy jest realny, ale występuje na innym etapie, niż go mierzymy. Oszczędzasz na pisaniu, a płacisz na przeglądzie, testowaniu i utrzymaniu. W firmwarze, gdzie przegląd wymaga rozumienia sprzętu, a test wymaga stanowiska pomiarowego, koszt tej drugiej strony bilansu jest wyższy niż gdziekolwiek indziej. Bilans wyjdzie na plus wyłącznie tam, gdzie proces został do tego przygotowany – i temu poświęcona jest reszta artykułu.

3. Co modele naprawdę potrafią na mikrokontrolerze

Rankingi modeli, które trafiają do prasy, opierają się na testach porównawczych w rodzaju HumanEval czy SWE-bench. Pierwszy sprawdza, czy model napisze funkcję w Pythonie na podstawie opisu w komentarzu. Drugi – czy naprawi zgłoszenie w repozytorium



Rysunek 4. Degradacja nie jest proporcjonalna do trudności zadania. Badanie 27 modeli na ośmiu scenariuszach z układem ESP32

projektu otwartego, także niemal zawsze pythonowym. Ani jeden, ani drugi nie mówi niczego o tym, co się stanie, gdy kod trafi na układ z trzydziestoma sześcioma kilobajtami pamięci RAM i przerwaniem, które musi zmieścić się w dwunastu mikrosekundach.

Od końca 2025 roku zaczęły się jednak pojawiać testy porównawcze budowane specjalnie dla systemów wbudowanych – i to one są dziś najciekawszym źródłem wiedzy o rzeczywistych możliwościach tych narzędzi.

Urwisko, nie zbocze

Najbardziej wymowne dane pochodzą z badania opublikowanego w lutym 2026 roku w czasopiśmie „Future Internet”. Autorzy przetestowali dwadzieścia siedem modeli na ośmiu scenariuszach o rosnącej złożoności, zbudowanych na układzie ESP32 z czujnikami środowiskowymi i czujnikiem ruchu – od zwykłego odczytu czujnika po pełną integrację z bazą danych w chmurze wraz z panelem wizualizacji. Każdy wynik oceniało niezależnie dwóch recenzentów, metodą hierarchicznej analizy problemu.

Czołówka wypadła bardzo dobrze: trzy najlepsze modele uzyskały oceny w przedziale od 0,984 do 0,910. Interesujące jest jednak nie to, kto wygrał, tylko kształt krzywej. **Dla zadań prostych kod, który się kompiluje, wygenerowało od dziewiętnastu do dwudziestu trzech modeli. Dla zadań złożonych – od trzech do pięciu.**

To nie jest łagodne zbocze. To urwisko. Pogorszenie nie następuje proporcjonalnie do trudności, lecz skokowo, w momencie, w którym zadanie przestaje być złożeniem znanych fragmentów, a zaczyna wymagać utrzymania spójności między kilkoma podsystemami naraz.

Praktyczna konsekwencja dla warsztatu jest taka: jeżeli twoje doświadczenie z asystentem opiera się na zadaniach z lewej strony tej krzywej – miganie diodą, odczyt czujnika, szkielek obsługi portu szeregowego – to twoja ocena jego przydatności jest systematycznie zawyżona względem tego, co dostaniesz przy integracji.

Trzy testy, które mierzą to, co trzeba

EmbedBench wraz z ramami badawczymi EmbedAgent, przyjęty na tegoroczną konferencję ICSE, jako pierwszy objął pełną ścieżkę: od projektu obwodu, przez kod, po przeniesienie rozwiązania na inną platformę sprzętową. Model dostaje opis zadania oraz schemat i ma napisać firmware. Trudność polega na utrzymaniu spójności między światem sygnałów, w którym moduły pracują równolegle i są połączone przepływem, a światem wywołań, w którym program biegnie sekwencyjnie po grafie funkcji. Modele wyszkolone na tym drugim przenoszą się do pierwszego gorzej, niż sugerują ich wyniki w testach ogólnych.

IoT-SkillsBench, opracowany na Duke University, obejmuje trzy platformy sprzętowe, dwadzieścia trzy peryferia

i czterdzieści dwa zadania w trzech poziomach trudności, sprawdzane wykonaniem na rzeczywistym sprzęcie. Najważniejszy jest jednak układ eksperymentu: każde zadanie uruchamiano w trzech konfiguracjach agenta – bez dodatkowych procedur, z procedurami wygenerowanymi przez sam model oraz z procedurami napisanymi przez człowieka-eksperta. **Wariant z procedurami eksperckimi wygrywa.** Jest to wynik o dużym znaczeniu dla naszego zawodu, ponieważ mówi, że wiedza inżynierska nie staje się zbędna – przenosi się o poziom wyżej, do warstwy instrukcji, którymi obudowuje się model. Wróćmy do tego tematu we fragmencie o kontekście.

Embedded Arena poszedł najdalej. Agent kolejnymi poprawkami modyfikuje firmware, oprogramowanie badawcze kompiluje je i wgrywa na płytkę, a ocena pochodzi z pomiarów: czy obraz w ogóle daje się wdroić, ile pobiera prądu, ile zużywa energii, jak bardzo się nagrzewa. Jest to pierwszy test porównawczy, w którym „działa” znaczy „mieści się w budżecie mocy”, a nie „przechodzi asercję”.

Awarie, które nie wyglądają jak awarie

Odrębny wątek wnosi praca Roberta Morabita i Guanghai Wu, opublikowana w listopadzie 2025 roku w miesięczniku IEEE „Computer”. Autorzy zbudowali kompletny łańcuch przetwarzania dla uczenia maszynowego na urządzeniach wbudowanych, sterowany modelem językowym – przygotowanie danych, konwersja modelu, generowanie kodu wniosku na urządzeniu – i skatalogowali sposoby, w jakie ten łańcuch się psuł.

Najcenniejszy wniosek dotyczy kategorii, którą nazwali kodem rozsypującym się w czasie wykonania. Fragment kompiluje się bez ostrzeżenia, przechodzi sprawdzenie na swoim etapie, a rozbija dopiero etap kolejny

Czego HumanEval i SWE-bench nie mierzą

Pozycja modelu w rankingu ogólnym nie przekłada się na firmware, ponieważ testy ogólne nie zawierają:

- **sprzętu w pętli** – nie ma płytki, więc nie ma sposobu, by kod zawiódł dopiero na niej;
- **budżetu czasowego** – poprawność jest dwuwartościowa, nie istnieje pojęcie najgorszego przypadku ani nieprzekraczalnego terminu reakcji;
- **rejestrów i erraty** – interfejs programowy jest udokumentowany i stabilny, nie ma bitów zachowujących się inaczej, niż zapisano w nocie;
- **budżetu pamięci i energii** – nikt nie sprawdza rozmiaru sekcji ani poboru prądu;
- **współbieżności z przerwaniami** – brak procedur obsługi przerwań, brak sekcji krytycznych, brak wyścigów o dostęp do zmiennych;
- **konsekwencji długu** – zadanie kończy się w chwili, gdy testy świecą na zielono; nikt nie utrzymuje tego kodu przez dziesięć lat.

Wniosek praktyczny: przy wyborze modelu do pracy nad firmware sięgaj po testy z pomiarem sprzętowym, a nie po tabele wyników SWE-bench.

– i to w sposób, którego typowe procedury walidacyjne nie wykrywają. Autorzy pokazali też, że sam format zapytania wpływa nie tylko na skuteczność, lecz również na charakterystykę popełnianych błędów. Oznacza to, że dwa zespoły używające tego samego modelu do tego samego zadania mogą otrzymywać systematycznie różne rodzaje usterek, zależnie od tego, jak sformułowały polecenie.

4. Anatomia typowego błędu w kodzie na mikrokontroler

Przejdźmy do konkretności. Poniższe kategorie nie są listą teoretycznych zagrożeń, lecz tym, co rzeczywistość powraca w kodzie generowanym dla mikrokontrolerów – uporządkowanym według tego, jak trudno dany błąd wykryć.

Rejestry i bity konfiguracyjne. Model swobodnie miesza rodziny układów. Bity konfiguracyjne z układu PIC12F629 pojawiają się w projekcie na PIC16F, wywołania biblioteki sprzętowej charakterystyczne dla STM32F4 trafiają do projektu na STM32Go, gdzie odpowiednie peryferium ma inną strukturę rejestrów albo nie istnieje. Zdarzają się również pola struktur, których nie ma w żadnej rodzinie – model konstruuje je przez analogię do nazw, które widział w innych projektach.

Zegary i budżet czasowy. Klasykiem jest funkcja opóźniająca wywoływana w pętli głównej tam, gdzie potrzebny jest licznik sprzętowy. Kod działa. Rdzeń przez cały ten czas kręci się w pętli, więc o trybach niskiego poboru mocy można zapomnieć, a każde przerwanie wchodzące w środek opóźnienia rozjeżdża okresowość. Model nie analizuje najgorszego przypadku, ponieważ nie ma pojęcia o istnieniu takiej kategorii.

Współbieżność. Brak kwalifikatora volatile przy zmiennej dzielonej z procedurą obsługi przerwania sprawia, że kompilator przy włączonej optymalizacji trzyma ją w rejestrze, a pętla nigdy nie zauważy zmiany. Sekcja krytyczna bywa pominięta przy dostępie do struktury modyfikowanej z przerwania. Na rdzeniu Cortex-M0+ dochodzi brak instrukcji dostępu wyłącznego, oznaczanych LDREX i STREX, którymi rdzenie

M3 i wyższe realizują niepodzielne sekwencje odczytu, modyfikacji i zapisu. Warto to rozróżnić, ponieważ bywa mylone: sam odczyt albo sam zapis poprawnie wyrównanej zmiennej trzydziestodwubitowej jest na M0+ niepodzielny. Niepodzielne nie jest natomiast zwiększenie licznika ani ustawienie pojedynczego bitu w rejestrze, ponieważ każda z tych operacji rozkłada się na odczyt, zmianę i zapis, a między nie może wejść przerwanie. Model przenosi tu rozwiązanie z kodu pisanego dla rdzeni M4 i wyższych, gdzie mechanizm dostępu wyłącznego istnieje.

Pamięć. Alokacja dynamiczna w firmwarze, bufora zakładane na stosie wewnątrz procedury przerwania, formatowanie liczb zmiennoprzecinkowych na rdzeniu bez sprzętowej jednostki zmiennoprzecinkowej. Ten ostatni przypadek to kilkanaście kilobajtów kodu doklejonych przez konsolidator i emulacja programowa tam, gdzie wystarczyłaby arytmetyka stałoprzecinkowa.

Peryferia i świat fizyczny. Pominięty rezystor podciągający, brak eliminacji drgań styków, zła kolejność uruchamiania – konfiguracja portu przed włączeniem zegara peryferium jest tu najczęstsza. Osobną kategorię stanowi stan aktywny wyprowadzenia: dioda podłączona tak, że wyprowadzenie zwiera prąd do masy, świeci przy stanie niskim. Jeżeli tej informacji nie ma w kontekście, logika wyjdzie odwrotna, chociaż kod będzie wyglądał bez zarzutu.

Zmyślone interfejsy i pakiety. Nieistniejące funkcje bibliotek sprzętowych, wywołania z wersji zestawu narzędzi sprzed trzech lat, wymyślone nazwy komponentów i bibliotek. Ta ostatnia kategoria ma wymiar znacznie poważniejszy niż błąd kompilacji i poświęcamy jej osobny fragment punktu o bezpieczeństwie.

Jak to wygląda w kodzie

Poniższy listing powstał w odpowiedzi na polecenie typowe dla codziennej pracy: „Napisz funkcję w języku C dla układu STM32G071, odczytującą temperaturę z czujnika TMP102 przez magistralę I²C, z obsługą błędów”. Kod jest sformatowany, opatrzone komentarzami i kompiluje się bez ostrzeżeń przy domyślnych ustawieniach. Zawiera cztery wady.

```
/* Listing 1. Odczyt czujnika TMP102 – wersja wygenerowana */
#define TMP102_ADDR 0x48 /* (1) */
#define TMP102_REG_T 0x00

float TMP102_ReadTemperature(void)
{
    uint8_t reg = TMP102_REG_T;
    uint8_t buf[2];

    HAL_I2C_Master_Transmit(&hi2c1, TMP102_ADDR, &reg, 1,
                           HAL_MAX_DELAY); /* (2) */
    HAL_I2C_Master_Receive (&hi2c1, TMP102_ADDR, buf, 2,
                           HAL_MAX_DELAY); /* (2) */

    int16_t raw = (int16_t)((buf[0] << 8) | buf[1]); /* (3) */
    return raw * 0.0625f; /* (4) */
}
```

```

/* Listing 2. Ta sama funkcja po poprawce */
#define TMP102_ADDR_7B 0x48
#define TMP102_ADDR_HAL (TMP102_ADDR_7B << 1) /* HAL: adres 8-bitowy */
#define TMP102_REG_T 0x00
#define TMP102_TIMEOUT 50u /* ms – patrz budżet czasowy */

/* Wynik w setnych częściach stopnia: 2312 = 23,12 stopnia Celsjusza */
bool TMP102_ReadTemperature(int16_t *out_centi_c)
{
    uint8_t reg = TMP102_REG_T;
    uint8_t buf[2];

    if (HAL_I2C_Master_Transmit(&hi2c1, TMP102_ADDR_HAL, &reg, 1,
                                TMP102_TIMEOUT) != HAL_OK) {
        return false;
    }
    if (HAL_I2C_Master_Receive(&hi2c1, TMP102_ADDR_HAL, buf, 2,
                                TMP102_TIMEOUT) != HAL_OK) {
        return false;
    }

    /* Wyrównanie do lewej stawia bit znaku na pozycji 15, czyli tam,
    gdzie oczekuje go int16_t. Znak wychodzi poprawnie sam. */
    int16_t raw = (int16_t)((((uint16_t)buf[0] << 8) | buf[1]));

    /* To samo wyrównanie czyni wartość 16 razy większą, stąd 625/1600
    zamiast 0,0625. Mnożenie jawnie 32-bitowe, bez cichej promocji. */
    *out_centi_c = (int16_t)((((int32_t)raw * 625)/1600));
    return true;
}

```

(1) **Adres nieprzesunięty.** Czujnik TMP102 ma adres siedmiobitowy 0x48, natomiast interfejs biblioteki sprzętowej oczekuje adresu ośmiobitowego, czyli przesuniętego w lewo o jeden bit. Transmisja idzie pod adres 0x24 i nikt nie odpowiada. Kompilator tego nie zobaczy, analiza statyczna również – jest to poprawny kod z błędną wartością.

(2) **Brak limitu czasu.** Stała HAL_MAX_DELAY oznacza oczekiwanie bez końca. Przy zwarcu linii danych do masy albo przy niewlutowanym czujniku funkcja nie wraca nigdy, a wraz z nią zawiesza się cały wątek. Jest to dokładnie ta kategoria konstrukcji maskujących błędy, której udział wzrósł o czterdzieści siedem procent. Dodatkowo status obu wywołań jest ignorowany, więc

informacja o niepowodzeniu przepada, nawet gdyby limit czasu istniał.

(3) **Zły współczynnik skali.** Czujnik TMP102 podaje wynik dwunastobitowy, wyrównany do lewej w słowie szesnastobitowym: bity znaczące zajmują pozycje od piętnastej do czwartej, a cztery najmłodsze są zerami. Ma to jeden dobry skutek uboczny, który warto odnotować, ponieważ bywa źródłem nieporozumień. Bit znaku trafia dokładnie na pozycję piętnastą, czyli tam, gdzie oczekuje go typ int16_t, więc konwersja wykonana w tym listingu odtwarza wartości ujemne poprawnie i niczego nie trzeba do niej dopisywać. Wyrównanie sprawia natomiast, że liczba jest szesnastokrotnie większa od odczytu czujnika, podczas gdy mnożnik podany w następnej linii odnosi

Tabela 1. Typ błędu a narzędzie wykrywające

Typ błędu	Komp.	Stat.	Host	Sprzęt	Osc.	Prąd	Człow.
Zmyślona nazwa funkcji lub rejestru	●	●					●
Pomieszone rodziny mikrokontrolerów	●	●		●			●
Zły adres peryferium na magistrali				●	●		●
Brak limitu czasu w pętli oczekiwania		●		●			●
Brak volatile przy zmiennej z przerwaniami		●		●			●
Odczyt, modyfikacja i zapis bez ochrony na M0+		●		●			●
Zła konwersja lub skalowanie wyniku			●	●			●
Przekroczenie budżetu czasowego				●	●		
Przepełnienie stosu		●		●			
Emulacja zmiennoprzecinkowa, rozrost obrazu		●					●
Nadmierny pobór prądu						●	●
Zignorowany błąd układu z erraty				●	●		●
Powielenie kodu, rozjazd architektury		●					●

się do wartości nieprzesuniętej. Wynik jest więc szesnaście razy zawyżony w całym zakresie, także dla temperatur ujemnych – i to właśnie czyni ten błąd trudniejszym do wychwycenia, niż byłby przy uszkodzonym znaku. Przebieg rośnie i opada tak, jak powinien, tylko w złej skali. Do naprawy wystarczyłby mnożnik szesnaście razy mniejszy.

(4) Typ zwracany i brak kanału zgłaszania błędu. Typ `float` na rdzeniu Cortex-M0+ bez jednostki zmiennoprzecinkowej oznacza emulację programową. Poważniejsze jest jednak to, że funkcja nie ma jak zasygnalizować niepowodzenia – typowym „rozwiązaniem” generowanym w takich przypadkach jest podawanie wartości umownej w rodzaju `-999.0f`, którą kod wywołujący prędzej czy później potraktuje jako rzeczywisty pomiar.

Warto zwrócić uwagę, którą wadę czym wykryto. Punkt czwarty wyłapuje przegląd kodu oraz raport konsolidatora. Punkt trzeci – test jednostkowy uruchomiony na komputerze, w którym znanej wartości surowej z rejestru odpowiada znana temperatura. Punkt drugi – dopiero test na sprzęcie z celowo wymuszoną awarią magistrali. Punkt pierwszy – wyłącznie analizator stanów logicznych albo działający czujnik. **Żadnej z tych czterech wad nie zgłosi kompilator.**

Objaśnienie kolumn tabeli 1: **Komp.** – kompilator; **Stat.** – analiza statyczna; **Host** – test jednostkowy na komputerze; **Sprzęt** – test na rzeczywistym układzie; **Osc.** – oscyloskop lub analizator stanów logicznych; **Prąd** – pomiar poboru prądu; **Człow.** – przegląd wykonany przez człowieka. ● wykrywa niezawodnie ○ wykrywa częściowo lub zależnie od konfiguracji

Najważniejsze w tej tabeli są puste pola. Dwie ostatnie kategorie – zignorowany błąd układu oraz rozjazd architektoniczny – nie mają ani jednego pełnego oznaczenia poza ostatnią kolumną. Są to dokładnie te błędy, których żadne narzędzie automatyczne nie zgłosi, a które w kodzie tworzonym maszynowo występują częściej niż w pisanym ręcznie. I to jest praktyczne uzasadnienie tezy, do której wracamy w rozdziale warsztatowym: **generator nie może być jedynym walidatorem.**

5. Kontekst jest wszystkim: protokół MCP i procedury w praktyce

Wróćmy do diagnozy postawionej w pierwszym punkcie. Model językowy zawodzi w firmwarze nie dlatego, że jest zbyt słaby, lecz dlatego, że pracuje na tej części opisu systemu, którą mu pokazano – a w naszej branży pokazuje mu się zwykle najmniej istotny fragment, czyli sam kod źródłowy.

Zastanówmy się, gdzie w projekcie na mikrokontroler rzeczywiście znajduje się wiedza potrzebna do napisania choćby jednej poprawnej funkcji obsługi peryferium.

Część siedzi w nocie katalogowej układu: mapa rejestrów, zależności czasowe, warunki elektryczne. Część w erracie, czyli w wykazie błędów samego układu – dokumencie, którego istnienia model często nawet nie podejrzewa. Część na schemacie ideowym, bo to schemat mówi, że dioda jest podłączona katodą do wyprowadzenia i świeci przy stanie niskim. Część w drzewie zegarów wygenerowanym przez konfigurator, bo od niego zależy, czy dzielnik magistrali da żadaną częstotliwość taktowania. Część w pliku opisu sprzętu `devicetree` albo w konfiguracji `Kconfig`. Część w raporcie konsolidatora, który pokazuje, ile pamięci zostało. Część w logu kompilatora. A część – i bywa to najważniejsza część – w przebiegu zarejestrowanym analizatorem stanów logicznych na stanowisku, przy poprzedniej rewizji płytki.

Każde z tych źródeł jest osobnym silosem, w innym formacie, często w innym systemie. Inżynier scala je w głowie i to scalanie jest w gruncie rzeczy istotą zawodu. Model dostaje z tego zwykle jedno źródło – kod – a resztę uzupełnia z pamięci wyuczonej na cudzych projektach. Stąd biorą się pomieszane rodziny układów, nieistniejące pola struktury i odwrócona logika diody.

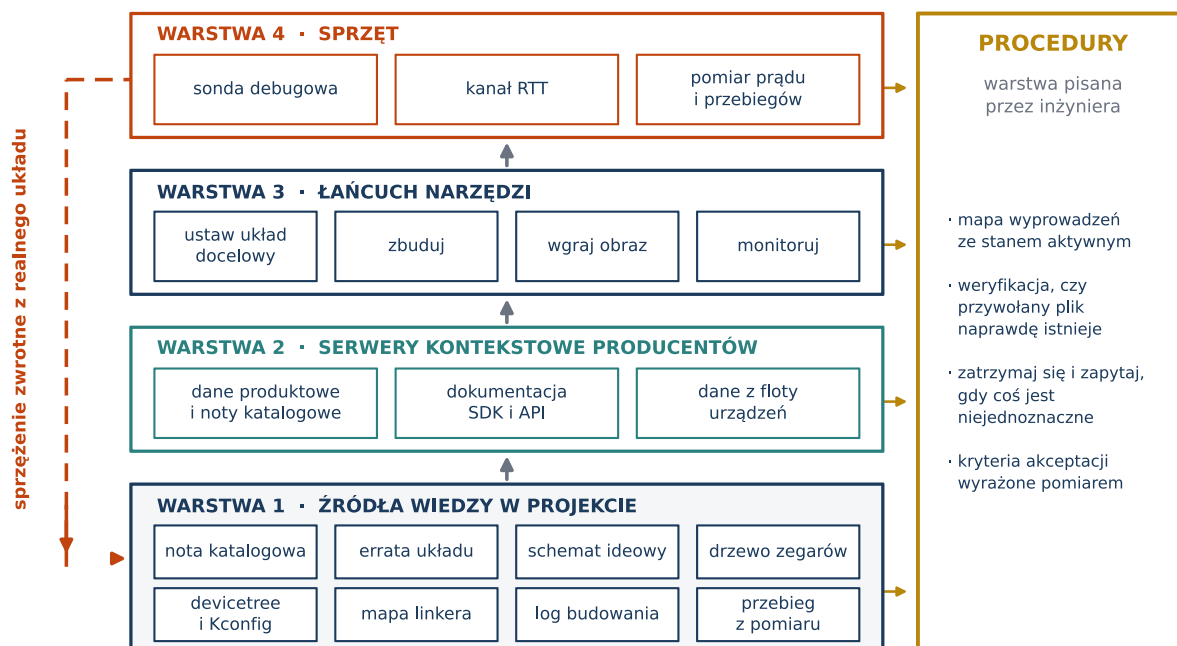
Protokół kontekstowy, znany pod skrótem MCP, jest próbą rozwiązania dokładnie tego problemu. W uproszczeniu jest to standard, w którym zewnętrzne źródło danych albo narzędzie wystawia asystentowi zestaw funkcji do wywołania wraz z opisem, co każda z nich robi i jakich argumentów oczekuje. Asystent nie musi wiedzieć z góry, co potrafi dany serwer – pyta go o to i otrzymuje listę. Z punktu widzenia inżyniera efekt jest taki, że **model przestaje zgadywać, a zaczyna sprawdzać.**

Dla oprogramowania wbudowanego ma to znaczenie większe niż dla jakiegokolwiek innej dziedziny programowania, właśnie dlatego, że nasze silosy wiedzy są głębsze i liczniejsze. W praktyce ekosystem, który wyrósł wokół tego protokołu przez ostatnie kilkanaście miesięcy, układa się w cztery wyraźne warstwy.

Warstwa pierwsza: dane producenta układu

Firma Microchip uruchomiła swój serwer kontekstowy w listopadzie 2025 roku i udostępniła go bezpłatnie, bez wymogu uwierzytelniania. Asystent może przez niego pobrać zweryfikowane dane o produktach: specyfikacje, odnośniki do not katalogowych i przewodników, typ oraz wymiary obudowy, liczbę wyprowadzeń, a także informacje o zgodności – dyrektywa RoHS, klasyfikacja eksportowa ECCN, poziom wrażliwości na wilgoć, kraj pochodzenia. Co ciekawe, ten sam serwer podaje również stany magazynowe, ceny i terminy realizacji, co czyni z niego narzędzie przydatne nie tylko przy pisaniu kodu, lecz także na etapie doboru układu, kiedy pytanie „czy ten

Model, któremu pokazano wyłącznie kod źródłowy, resztę uzupełnia z cudzych projektów



Rysunek 5. Cztery warstwy kontekstu w projekcie na mikrokontroler. Model, któremu pokazano wyłącznie kod źródłowy, resztę uzupełnia z cudzych projektów

mikrokontroler będzie dostępny za pół roku” bywa ważniejsze od pytania o jego peryferia.

Firma Nordic Semiconductor poszła w innym kierunku i wystawiła asystentowi dokumentację zestawu nRF Connect SDK wraz z opisem interfejsów programowych i konfiguracjami urządzeń, a dodatkowo – dane z floty urządzeń pracujących w usłudze chmurowej nRF Cloud. To drugie jest jakościowo nowe: model przestaje pracować wyłącznie na dokumentacji, a zaczyna mieć wgląd w to, jak zachowują się rzeczywiste egzemplarze w terenie.

Warstwa druga: łańcuch narzędzi

W kwietniu 2026 roku, wraz z wydaniem zestawu ESP-IDF w wersji 6.0, firma Espressif dołączyła do niego lokalny serwer kontekstowy obsługujący same narzędzia budujące. Asystent może przez niego ustawić układ docelowy, uruchomić kompilację, wgrać obraz na płytkę i sprawdzić stan projektu – a wszystko to z poziomu klienta pokroju Cursora czy Claude Code, bez opuszczania okna rozmowy.

Jest to moment, w którym warto się zatrzymać, ponieważ wtedy przebiega granica opisana w ramce A. Dopóki asystent tylko proponuje kod, człowiek pozostaje operatorem: kopiuje, kompiluje, ocenia. Od chwili, gdy asystent sam wywołuje polecenie budowania i sam czyta jego wynik, człowiek staje się recenzentem cudzej pracy – i to pracy, której pośrednie etapy zobaczy tylko wtedy, gdy celowo do nich zajrzy. Zysk jest realny, bo pętla „popraw i sprawdź” skraca się z minut do sekund. Koszt również jest realny i polega na tym, że pomiędzy jedną a drugą decyzją człowieka przybywa teraz znacznie więcej zmian.

Warstwa trzecia: sprzęt

Najdalej idą narzędzia dające asystentowi dostęp do sondy debugowej. Otwarty projekt zbudowany na bibliotece probe-rs udostępnia dwadzieścia dwa narzędzia obsługujące układy ARM Cortex-M oraz RISC-V przez sondy J-Link, ST-Link i pokrewne: zarządzanie sondą, odczyt i zapis pamięci, sterowanie wykonaniem programu, pułapki, operacje na pamięci Flash, komunikację przez kanał RTT i zarządzanie sesją. Istnieje też podobny serwer dla układów Texas Instruments z rodziny CC26xx oraz Nordic nRF52 i nRF91, obsługujący dodatkowo przechwytywanie transmisji z portu szeregowego.

Zastosowanie praktyczne wygląda tak, że inżynier opisuje objaw słowami, a agent samodzielnie wykrywa układ, kasuje pamięć, wgrywa obraz, zeruje cel i pokazuje logi rozruchu. Przy uruchamianiu nowej płytki oszczędność czasu bywa znaczna.

Trzeba jednak wyraźnie powiedzieć, czego to dotyczy. Kasowanie całej pamięci użytkownika jest operacją nieodwracalną. Niektóre układy zawierają program rozruchowy, który po wgraniu trwale wyłącza interfejs JTAG, a jedyną drogą odzyskania jest właśnie pełne kasowanie – o ile ten program na nie pozwoli. Agent wykonujący łańcuch operacji na podstawie zdania w języku naturalnym nie odróżnia płytki prototypowej od egzemplarza z serii produkcyjnej stojącego obok na biurku. Fizyczne oddzielenie stanowiska doświadczalnego od czegokolwiek, co pojedzie do klienta, przestaje być kwestią porządku i staje się zabezpieczeniem.

Warstwa czwarta: procedury, czyli tam, gdzie przenosi się warsztat

Najciekawsza warstwa nie polega na dostępie do danych ani do sprzętu, lecz na tym, jak opisać agentowi sposób postępowania. Przypomnijmy wynik z testu IoT-SkillsBench z poprzedniego punktu: procedury napisane przez człowieka-eksperta dają lepsze rezultaty niż procedury wygenerowane przez sam model. Zobaczmy, jak taka procedura wygląda w praktyce.

Jacob Beningo opisał w lipcu 2026 roku warsztat pracy nad plikiem opisu sprzętu w systemie Zephyr. Kuszące jest, pisać, otworzyć okno asystenta i poprosić o wygenerowanie takiego pliku dla nowej płytki. Byłby to błąd, ponieważ model nie jest deterministyczny – chyba że obuduje się go procedurą, która sama sprawdza własne wyniki i sama weryfikuje, czy model nie zmyśla.

Praktyczne elementy tej procedury są proste i dają się przenieść na dowolną platformę. Po pierwsze, mapa wyprowadzeń jako jedyne źródło prawdy, z obowiązkową kolumną stanu aktywnego – ponieważ to jest kolumna, którą ludzie pomijają i która najczęściej wywraca logikę. Po drugie, obowiązek sprawdzenia, że przywołany przez model plik opisu istnieje w drzewie projektu, zanim cokolwiek zostanie zapisane. Po trzecie, reguła zachowania przy niejednoznaczności: wiersz oznaczony jako „do ustalenia” albo dwa wiersze roszczące sobie prawo do tego samego wyprowadzenia stanowią pytanie do konstruktora sprzętu, a nie sytuację, w której agent po cichu wybiera jedną z możliwości.

Podobnie działa zestaw czterech reguł opublikowany przez Espressif jako ograniczenia dla projektów zephyrowych: wykorzystaj maksymalnie sam system operacyjny, zamiast budować własną warstwę pośrednią; wykorzystaj jego usługi – sieć, zarządzanie łącznością, stos Bluetooth, ustawienia, rejestrowanie zdarzeń, kolejki zadań – zamiast pisać ich zamienniki; wykorzystaj moduły z drzewa

dystrybucji zamiast prywatnych kopii; zintegruj się zgodnie z przyjętą konwencją, czyli przez przestrzeń roboczą west, devicetree i plik konfiguracyjny projektu. Autorzy zwracają uwagę, że są to znakomite ograniczenia zapytania, ponieważ mówią modelowi, jak wygląda dobra robota, zanim ten zdąży wymyślić obcą strukturę projektu.

Na tym właśnie polega przeniesienie warsztatu o poziom wyżej. Wiedza inżynierska nie znika i nie zostaje zastąpiona – zmienia postać z „umiem to napisać” na „umiem opisać, jak to ma być napisane i po czym poznać, że wyszło źle”.

6. Bezpieczeństwo: „kompiluje się” to za mało

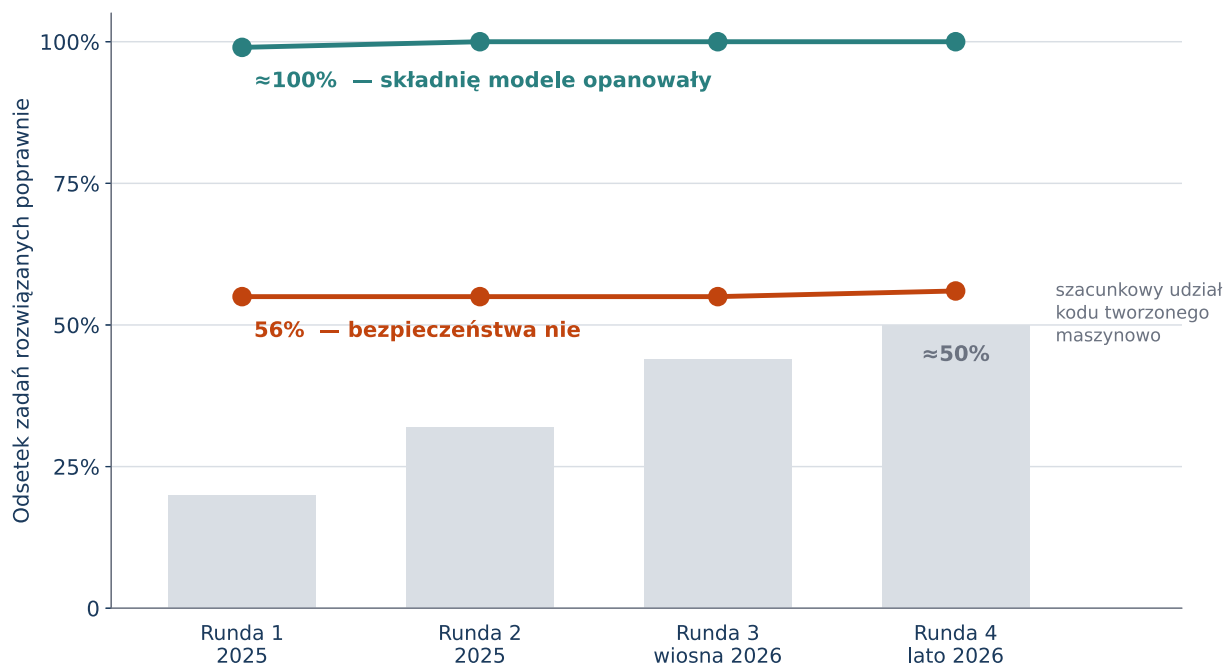
Wróćmy do liczby z początku artykułu i rozłóżmy ją na części, ponieważ sposób jej powstania ma znaczenie dla tego, jak ją czytać.

Firma Veracode prowadzi swój pomiar od 2025 roku i przebadła w sumie ponad sto pięćdziesiąt modeli językowych. Metoda jest prosta i przez to mocna: model dostaje osiemdziesiąt zadań polegających na uzupełnieniu funkcji zgodnie z komentarzem opisującym pożądane działanie. Zadania dobrano tak, że każde da się rozwiązać zarówno w sposób bezpieczny, jak i podatny. Nie ma żadnej podpowiedzi sugerującej, że bezpieczeństwo ma znaczenie – ponieważ w codziennej pracy takiej podpowiedzi też zwykle nie ma. Wynik przepuszcza się następnie przez analizator kodu, który sprawdza cztery ściśle określone klasy podatności z katalogu CWE, wybrane spośród listy OWASP Top 10. Są to: wstrzyknięcie zapytania do bazy danych, czyli SQL injection (CWE-89); wstrzyknięcie skryptu do strony wyświetlanej użytkownikowi, czyli cross-site scripting (CWE-80); wstrzyknięcie spreparowanego wpisu do dziennika zdarzeń, czyli log injection (CWE-117); oraz użycie algorytmu kryptograficznego uznanego za niebezpieczny, na przykład przestarzałego

Minimalny pakiet kontekstu dla agenta pracującego nad firmware

Zestaw materiałów, bez których nie warto zaczynać. Kolejność odpowiada temu, jak często ich brak wywraca wynik.

- 1. Mapa wyprowadzeń ze stanem aktywnym** – sygnał, kierunek, stan aktywny, podciągnięcie, uwaga o funkcji alternatywnej. Wiersz niejednoznaczny jest pytaniem do konstruktora, nie zaś domyślnym wyborem agenta.
- 2. Rewizja płytki i pełne oznaczenie układu** – łącznie z wariantem obudowy i wielkością pamięci, ponieważ od niej zależy mapa pamięci.
- 3. Wyciągi z noty katalogowej i pełna errata układu** – mapa rejestrów peryferiów użytych w zadaniu oraz errata, czyli publikowany przez producenta wykaz miejsc, w których wyprodukowany układ zachowuje się inaczej, niż opisuje to nota katalogowa. Nie chodzi o błędy w kodzie, tylko o błędy w samym układzie scalonym, wykryte już po uruchomieniu produkcji i naprawiane dopiero w kolejnej rewizji. Typowe pozycje to zablokowanie magistrali I²C po przerwaniu transferze, przekłamany odczyt przetwornika analogowo-cyfrowego przy najkrótszym czasie próbkowania albo przerwanie zgłaszane raz zamiast dwa razy. Producent podaje przy każdej pozycji sposób obejścia problemu – i tego obejścia model nie zna, ponieważ nie ma go w kodzie, na którym się uczył.
- 4. Drzewo zegarów** – źródło, dzielniki, częstotliwości magistral, tryby oszczędzania energii dopuszczone w tym projekcie.
- 5. Mapa pamięci i bieżący stan jej wykorzystania** – ile zostało w pamięci programu i w pamięci RAM, gdzie leżą bufor statyczne, jaki jest rozmiar stosu.
- 6. Budżet czasowy** – najgorszy dopuszczalny czas obsługi przerwania, okresy zadań, priorytety.
- 7. Wersje łańcucha narzędzi** – kompilator, zestaw SDK, biblioteki, podane numerami wersji, a nie samymi nazwami.
- 8. Kryteria akceptacji** – wyrażone jako zachowanie mierzalne: przebieg na wyprowadzeniu, zgodność z protokołem, stan rejestrów po uruchomieniu, pobór prądu w trybie uśpienia.
- 9. Plan wycofania** – co zrobić, gdy zmiana okaże się zła, i skąd wziąć obraz poprzedni.



Rysunek 6. Ponad 100 modeli, 80 zadań, cztery rundy pomiaru. Odsetek się nie zmienia – zmienia się podstawa, do której się odnosi

funkcji skrótu albo klucza wpisanego na stałe w kodzie (CWE-327).

Rezultat z lata 2026 roku: średnia zdawalność pięćdziesiąt sześć procent, wobec pięćdziesięciu pięciu w pierwszym pomiarze. Najlepszy model w zestawie osiągnął sześćdziesiąt osiem procent, a więc przewraca się na blisko jednym zadaniu na trzy. Sześć z jedenastu badanych modeli mieści się w przedziale od pięćdziesięciu do pięćdziesięciu trzech procent. Modele budowane specjalnie do programowania wypadają średnio na pięćdziesięciu jeden procentach, czyli nie lepiej od ogólnych.

Kluczowa jest przy tym płaskość tej krzywej. Przez cztery lata zdolność modeli do rozwiązywania zadań programistycznych rosła w tempie, które wszyscy obserwujemy. Zdolność do rozwiązywania ich bezpiecznie nie drgnęła. Autorzy raportu wskazują przyczynę, która brzmi rozsądnie: modele uczą się na publicznych repozytoriach i serwisach z odpowiedziami, a ten zbiór odzwierciedla historyczne praktyki programistyczne razem z historycznymi błędami. Wzorce niebezpieczne, powszechne w 2020 roku, nadal są w tych danych i nadal stanowią większość przykładów danej konstrukcji. Nie jest to więc problem, który rozwiąże następna wersja modelu – jest to problem zbioru, na którym modele się uczą.

Średnia zdawalność zaciera przy tym rzecz najciekawszą, ponieważ te cztery klasy wypadają skrajnie różnie. Kryptografię modele obsługują dobrze, ze zdawalnością osiemdziesięciu siedmiu procent, a wstrzyknięcie zapytania do bazy danych – osiemdziesięciu trzech. Przy wstrzyknięciu skryptu do strony zdawalność spada do piętnastu procent, a przy wstrzyknięciu wpisu do dziennika zdarzeń do dwunastu. Wyjaśnienie tej różnicy jest dla nas ważniejsze niż same

liczby. Dobór bezpiecznego algorytmu szyfrowania czy użycie zapytania parametryzowanego to wzorce zamknięte, opisane w tysiącach przykładów i dające się rozpoznać po samym kształcie kodu. Natomiast rozstrzygnięcie, czy dana zmienna zawiera dane pochodzące z zewnątrz i wymaga oczyszczenia, wymaga prześledzenia, skąd ta wartość przyszła – przez kolejne wywołania funkcji, często między plikami. Modele tego nie robią, nawet dysponując dużym oknem kontekstu, i to jest właściwa granica ich możliwości.

Jedno zastrzeżenie trzeba postawić wprost, ponieważ dotyczy przenoszalności tych wyników na nasz teren. Test obejmuje cztery języki: Java, JavaScript, Python i C#. Języka C w nim nie ma, a więc nie są to pomiary wykonane na kodzie dla mikrokontrolerów. Przenosi się jednak sam mechanizm, i to w sposób, który dotyka firmware'u bezpośrednio. Nierozpoznanie, które dane pochodzą z zewnątrz, w aplikacji webowej kończy się wstrzykniętym skryptem, a w firmware – parserem ramki z portu szeregowego albo z łącza radiowego, który przyjmuje długość pola bez sprawdzenia i kopiuje ją do bufora o stałym rozmiarze. Jest to ta sama luka w rozumowaniu modelu, tyle że jej skutkiem nie jest wykradziona sesja, lecz nadpisany stos. Podobnie przekłada się kategoria kryptograficzna: klucz wpisany na stałe w kodzie jest tym samym błędem niezależnie od tego, czy trafi do serwisu internetowego, czy do pamięci programu w urządzeniu, które przepracuje u klienta dziesięć lat.

Dlaczego trafia to w firmware mocniej

Skoro języka C w tamtym teście nie ma, sięgnijmy po badanie, które go obejmuje – bo wypada w nim gorzej niż języki wyższego poziomu. W jednym z wcześniejszych, szeroko cytowanych eksperymentów przeanalizowano 1689

programów wygenerowanych przez asystenta i podatności znaleziono w około czterdziestu procentach z nich, przy czym dla kodu w języku C odsetek sięgał mniej więcej pięćdziesięciu procent, wobec około trzydziestu dziećciu dla Pythona. Powód jest oczywisty dla każdego, kto pisze na mikrokontrolery: w języku C odpowiedzialność za granice buforów, czas życia wskaźników i arytmetykę spoczywa wyłącznie na programiście, a modele odtworzają wzorce, w których ta odpowiedzialność bywała traktowana pobłażliwie.

Dochodzi do tego drugi efekt, opisany w tym samym nurcie badań: użytkownicy asystentów piszą kod mniej bezpieczny niż osoby pracujące bez nich, a jednocześnie **wyżej oceniają jego bezpieczeństwo**. Badacze nazwali to fałszywym poczuciem bezpieczeństwa, a mechanizm jest ten sam co przy rozjeździe percepcji i pomiaru w badaniu wydajności: kod tworzony maszynowo wygląda porządnie. Jest równo sformatowany, opatrzone komentarzami, nazwy zmiennych są sensowne. Wszystkie sygnały, po których człowiek zwyczajowo rozpoznaje kod pisany starannie, są obecne – i właśnie dlatego przegląd prześlizguje się po nim szybciej.

Efekt skali

Osobna analiza, przeprowadzona w firmach z listy Fortune 50, zestawiała dwie wielkości. Programiści wspomagani sztuczną inteligencją wprowadzają zmiany w tempie trzy do czterech razy wyższym niż pozostali. Zgłoszenia dotyczące bezpieczeństwa generują natomiast w tempie dziesięciokrotnie wyższym. Różnica między tymi mnożnikami jest właśnie tempem narastania długu – i narasta on szybciej, niż zespół zdąży go spłacać.

Znaczenie tego zestawienia jest takie, że niezmienny wskaźnik czterdziestu czterech procent oznacza dziś coś zupełnie innego niż dwa lata temu. Wtedy dotyczył

niewielkiego ułamka wprowadzanego kodu. Dziś, gdy szacuje się, że modele piszą mniej więcej połowę kodu trafiającego do repozytoriów, ten sam odsetek dotyczy wielokrotnie większej podstawy. Nie pogorszyła się jakość – zmieniła się skala.

Warto zestawzić to z niezależnym pomiarem z zupełnie innej strony: tegoroczny raport firmy Verizon o naruszeniach bezpieczeństwa wskazał podatności oprogramowania jako najczęstszy wektor wejścia, odpowiadający za trzydzieści jeden procent przypadków – po raz pierwszy przed skradzionymi poświadczeniami dostępu.

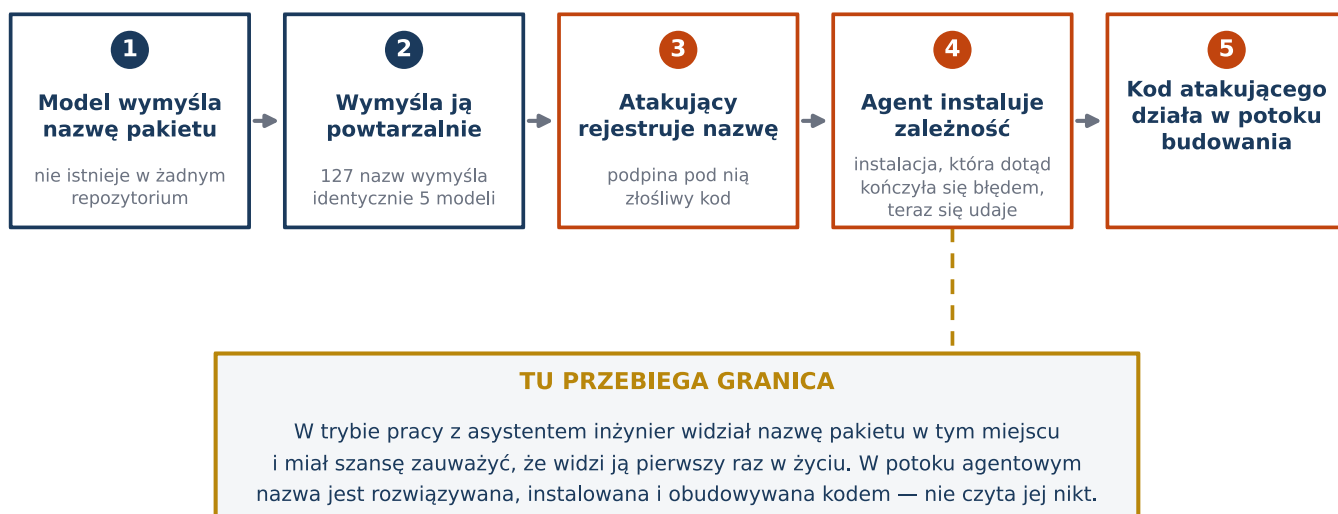
Pakiety, których nie ma

Osobnym i pod pewnym względem najbardziej podstawowym problemem są zmyślane nazwy zależności. Model, proszony o kod korzystający z biblioteki, potrafi wymyślić nazwę pakietu, który nie istnieje w żadnym repozytorium. Zjawisko opisano po raz pierwszy systematycznie w pracy przedstawionej na konferencji USENIX Security w 2025 roku, a nadany mu wtedy termin – *slopsquatting* – przyjął się w branży.

Mechanizm ataku jest krótki. Model wymyśla nazwę. Wymyśla ją powtarzalnie, przy różnych sformułowaniach zapytania. Atakujący rejestruje tę nazwę w publicznym repozytorium i podcina do niej złośliwy kod. Polecenie instalacji, które dotąd kończyło się błędem, zaczyna kończyć się powodzeniem – i uruchamia cudzy kod na maszynie inżyniera albo w piaskownicy agenta.

Powtórzenie tego badania na pięciu modelach wydanych między październikiem 2025 a marcem 2026 roku, obejmujące prawie dwieście tysięcy par zapytań weryfikowanych wobec list pakietów w repozytoriach PyPI i npm, dało wskaźniki zmyśleń od 4,62 do 6,10 procent. Jest to wyraźnie mniej niż wcześniej, a różnice między modelami się skurczyły – ale zagrożenie nie zniknęło.

Wskaźnik halucynacji nazw na modelach z 2026 roku: od 4,62 do 6,10 procent



Rysunek 7. Łańcuch ataku przez wymyśloną nazwę pakietu. Wskaźnik halucynacji nazw na modelach z 2026 roku: od 4,62 do 6,10 procent

Bramki między agentem a menedżerem pakietów

Żadna z tych kontroli nie działa samodzielnie. Sensowne jest ustawienie ich szeregowo.

- **Lista dozwolonych repozytoriów** – agent pobiera zależności wyłącznie z lustra kontrolowanego przez firmę, nigdy wprost z repozytorium publicznego.
- **Próg wieku i popularności pakietu** – nowa zależność poniżej ustalonego progu wymaga zatwierdzenia przez człowieka. Pakiet zarejestrowany w zeszłym tygodniu nie wchodzi do budowania automatycznie.
- **Plik blokady wersji jako element podlegający przeglądowi** – zmiana w nim jest zmianą wymagającą recenzji na równi ze zmianą w kodzie.
- **Ochrona przed pomyleniem repozytoriów** – jeżeli model wymyśli nazwę zbliżoną do waszego pakietu wewnętrznego, a nazwa ta jest wolna w repozytorium publicznym, mechanizm rozwiązywania zależności może sięgnąć na zewnątrz. Konfigurację trzeba sprawdzić celowo, a nie zakładać, że jest poprawna.
- **Analiza składu oprogramowania w potoku ciągłej integracji** – z wynikiem zapisywanym jako wykaz komponentów. Jest to ta sama praca, której i tak wymaga rozporządzenie omawiane w następnym rozdziale.

Zasada nadrzędna: bramka, której nigdy nie sprawdzono rzeczywistą próbą, jest założeniem, a nie zabezpieczeniem. Raz na kwartał warto celowo doprowadzić do zmyślenia nazwy, przepuścić ją przez cały potok i potwierdzić, że coś ją zatrzymało.

Autorzy zidentyfikowali sto dwadzieścia siedem nazw pakietów, które **wszystkie pięć badanych modeli wymyśla identycznie**. Nie jest to szum statystyczny, lecz gotowa lista celów.

Najistotniejsza zmiana dotyczy jednak nie samego wskaźnika, tylko tego, kto tę nazwę widzi. W trybie pracy z asystentem inżynier oglądał nazwę pakietu, zanim uruchomił instalację, i przynajmniej miał szansę zauważyć, że widzi ją pierwszy raz w życiu. W potoku agentowym nazwa jest rozwiązywana, instalowana, obudowywana kodem i zgłaszana jako gotowa zmiana. Nikt jej nie czyta. Udokumentowano już przypadek, w którym wymyślony pakiet rozszedł się przez ponad dwieście repozytoriów wraz z generowanymi maszynowo procedurami dla agentów, a jego pobrania napędzały nie osoby kopiujące kod, lecz agenty wykonujące własne wcześniejsze wyniki.

Dla nas dochodzi jeszcze wymiar prawny i poufnościowy. Blisko co piąty respondent badania RunSafe wskazał ryzyko związane z licencjami jako istotną obawę – modele uczą się na kodzie objętym różnymi warunkami i potrafią odtwarzać fragmenty bez wskazania źródła. Zdarza się również, że w danych uczących znalazły się poświadczenia dostępowe, co oznacza, że model może zaproponować w kodzie klucz albo hasło pochodzące z cudzego projektu.

7. Zgodność i odpowiedzialność: CRA, MISRA, bezpieczeństwo funkcjonalne

Wszystko, co opisaliśmy dotąd, było kwestią inżynierskiej rzetelności – a więc czegoś, co każdy zespół rozstrzyga po swojemu i na własne ryzyko. Od jesieni bieżącego roku przestaje tak być, przynajmniej w odniesieniu do wyrobów wprowadzanych na rynek Unii Europejskiej.

Kalendarz, który już się zaczął

Rozporządzenie Parlamentu Europejskiego i Rady (UE) 2024/2847, znane jako Cyber Resilience Act, weszło w życie 10 grudnia 2024 roku, ale jego obowiązki rozłożono

na blisko trzy lata. Warto rozróżnić trzy daty, ponieważ bywają mylone.

11 czerwca 2026 roku zaczął obowiązywać rozdział czwarty rozporządzenia, dotyczący jednostek notyfikowanych i organów krajowych. Państwa członkowskie miały do tego dnia wyznaczyć właściwe organy. Dla producenta oznacza to tyle, że aparat nadzoru już istnieje.

11 września 2026 roku wchodzi w życie obowiązek zgłaszania aktywnie wykorzystywanych podatności oraz poważnych incydentów bezpieczeństwa. Zegar jest bezlitosny: dwadzieścia cztery godziny na wczesne ostrzeżenie, siedemdziesiąt dwie godziny na zgłoszenie pełne, czterdzieści dni na raport końcowy po udostępnieniu poprawki, a w przypadku poważnych incydentów jeden miesiąc. Zgłoszenia idą przez wspólną platformę prowadzoną przez agencję ENISA, która instrukcje rejestracji opublikowała pod koniec lipca. Obowiązek dotyczy każdego wyrobu objętego zakresem rozporządzenia – również tego, który trafił na rynek wcześniej.

11 grudnia 2027 roku wchodzi pozostałe obowiązki: wymagania bezpieczeństwa na etapie projektowania, ocena zgodności, dokumentacja techniczna, oznakowanie CE, wykaz komponentów oprogramowania oraz obowiązek dostarczania aktualizacji przez cały deklarowany okres wsparcia. Po tej dacie wyrób bez oznakowania CE zgodnego z rozporządzeniem nie może być legalnie wprowadzony na rynek unijny.

Sankcje wynoszą do piętnastu milionów euro albo dwóch i pół procent światowego obrotu rocznego, w zależności od tego, która kwota jest wyższa.

Dlaczego dotyczy to kodu z asystenta

Na pierwszy rzut oka rozporządzenie nie mówi o sztucznej inteligencji ani słowa i można odnieść wrażenie, że są to dwa osobne tematy. Wrażenie jest mylne, a wiąże je jedna wielkość: dwadzieścia cztery godziny.

Kiedy okaże się, że w powszechnie używanej bibliotece jest podatność i że ktoś zaczął ją wykorzystywać, producent



Rysunek 8. Kalendarz Cyber Resilience Act i zegar zgłoszeniowy

ma dobę na wczesne ostrzeżenie. W tej dobie musi odpowiedzieć na pytanie, które brzmi trywialnie, a bywa niewykonalne: **które z moich wyrobów i w których wersjach oprogramowania zawierają ten komponent**. Nie da się na to odpowiedzieć przez przeszukiwanie repozytoriów i pytanie kolegów. Da się odpowiedzieć wyłącznie wtedy, gdy istnieje aktualny, czytelny maszynowo wykaz komponentów dla każdej wydanej wersji firmware'u, zestawiany na bieżąco z informacjami o podatnościach.

Formalnie obowiązek posiadania takiego wykazu – SBOM, ujęty w załączniku I, część II, punkt 1 rozporządzenia – wchodzi dopiero w grudniu 2027 roku. Praktycznie trzeba go mieć od września bieżącego roku, ponieważ bez niego obowiązek zgłoszeniowy jest nie do wykonania. Rozporządzenie nie narzuca konkretnego formatu, żądając jedynie formatu powszechnie stosowanego i czytelnego maszynowo, obejmującego co najmniej zależności najwyższego poziomu; w praktyce oznacza to CycloneDX albo SPDX. Wytyczne Komisji Europejskiej z 27 lipca bieżącego roku rozstrzygnęły przy tym kwestię, która budziła najwięcej wątpliwości: **obowiązek zgłoszenia obejmuje również podatności pochodzące z komponentów dostarczonych przez podmioty trzecie**. Nie ma tu przeniesienia odpowiedzialności na dostawcę biblioteki.

I tu wracamy do agenta. Agent działający w warstwie drugiej lub trzeciej z ramki A samodzielnie dociąga zależności. Robi to kompetentnie, szybko i zwykle sensownie – ale każda taka zależność jest komponentem, o którym nikt w firmie nie podjął świadomej decyzji, a który po grudniu 2027 roku musi znaleźć się w dokumentacji technicznej wyrobu i za który producent odpowiada już od września 2026. Zjawisko shadow AI, opisane w punkcie pierwszym jako ciekawostka organizacyjna, w tym świetle wygląda inaczej: jest to kod nieznanego pochodzenia,

ciągnący za sobą zależności nieznanego pochodzenia, w wyrobie objętym deklaracją zgodności.

Wniosek praktyczny jest jednoznaczny i wart wykonania niezależnie od tego, jak intensywnie korzystacie z asystentów: **generowanie wykazu komponentów musi być krokiem potoku budowania, a nie czynnością wykonywaną przed audytem**. Wykaz sporządzony ręcznie dezaktualizuje się, zanim skończycie go pisać, ponieważ każda aktualizacja firmware'u wymaga nowego. Narzędzia otwarte radzą sobie z tym w kilkadziesiąt sekund na jedno budowanie.

Stanowisko komitetu MISRA

Producenci oprogramowania dla motoryzacji i przemysłu mieli przez lata wygodną furtkę: wytyczne MISRA tradycyjnie łagodziły wymagania wobec kodu generowanego automatycznie, na przykład przez narzędzia modelowe pokroju Simulinka czy Stateflow. Część reguł doradczych uznawano wobec takiego kodu za niestosowane, co upraszczało wykazanie zgodności.

Komitet MISRA rozstrzygnął w wydaniu MISRA C:2025, że **kod pochodzący z modeli językowych z tej ulgi nie korzysta**. Uzasadnienie jest ściśle inżynierskie i warto je zapamiętać, ponieważ bywa przydatne w rozmowie z zarządem. Generator modelowy jest deterministyczny: z tego samego modelu przy tych samych ustawieniach otrzymuje się ten sam kod, a jego zachowanie jest przewidywalne i sprawdzalne raz dla całej klasy przypadków. Model językowy takich gwarancji nie daje. Wobec tego kod z niego pochodzący wymaga kontroli nie mniejszej, lecz **dotatkowej** względem kodu pisanego ręcznie: analiza statyczna, przegląd przez człowieka i testy jednostkowe są obowiązkowe, a nie zalecane.

Przy okazji warto odnotować dwa dokumenty towarzyszące, ponieważ wyznaczają kierunek dalszych prac. Nie

są to załączniki wewnątrz samej normy i w egzemplarzu MISRA C:2025 się ich nie znajdzie – konsorcjum publikuje je osobno, jako bezpłatne pliki PDF do pobrania ze swojej strony, i oba ukazały się w marcu 2025 roku wraz z nowym wydaniem wytycznych.

Pierwszy z nich, MISRA C:2025 Addendum 5, nosi tytuł „Coverage of MISRA C:2025 against the Common Weakness Enumeration” i zestawia wytyczne MISRA z katalogiem słabości oprogramowania CWE, prowadzonym przez korporację MITRE. Dokument przechodzi pozycją po pozycji przez ten katalog i dla każdej podaje, które reguły oraz dyrektywy MISRA ją pokrywają i jak mocno – od pełnego pokrycia przez regułę wprost jej poświęconą, przez pokrycie pośrednie, aż po jawne stwierdzenie, że danej słabości nie pokrywa żadna wytyczna. Interesuje nas w nim przede wszystkim blok dotyczący pamięci: przepełnienie bufora, zapis poza jego granicami, odczyt spoza zakresu tablicy, błędne obliczenie rozmiaru bufora. Jest to dokładnie ta klasa błędów, w której kod w języku C tworzony maszynowo wypada najgorzej, a jednocześnie ta, którą najtrudniej wykręcić przegłędem. Dokument ma też wyraźny cel polemiczny – grupa robocza MISRA odpowiada nim na powtarzaną w ostatnich latach tezę, że język C nie nadaje się do oprogramowania o wysokim poziomie integralności, i przedstawia mapowanie jako dowód, że przy odpowiednich rygorach nadaje się nadal. Adres: <https://misra.org.uk/app/uploads/2025/03/MISRA-C-2025-ADD5.pdf>

Drugi dokument, MISRA C:2025 Addendum 6, zatytułowany „Applicability of MISRA C:2025 to the Rust Programming Language”, ocenia, które z wytycznych napisanych dla języka C zachowują sens w języku Rust. Ocena ma postać macierzy z osobnymi kolumnami dla Rusta jako całości – obejmującego kod bezpieczny, bloki oznaczone jako niebezpieczne oraz wywołania bibliotek napisanych w innych językach – i dla samego podzbioru bezpiecznego. Wynik jest taki, że część wytycznych odpada, ponieważ język eliminuje problem u źródła, część zostaje w mocy, choć wymaga innego sposobu spełnienia, a część pozostaje w pełni aktualna. We wstępie autorzy stawiają sprawę wprost: krąży przekonanie, że samo przejście na Rusta rozwiązuje wszystkie problemy, a tak nie jest. Nie są to jeszcze reguły przeznaczone dla Rusta – konsorcjum zapowiada je jako przedmiot dalszych prac, we współpracy z grupą roboczą po stronie tego języka. Adres: <https://misra.org.uk/app/uploads/2025/03/MISRA-C-2025-ADD6.pdf>

Otwarte pytanie o kwalifikację narzędzia

Norma ISO 26262 w części ósmej wymaga, by każde narzędzie użyte w projekcie związanym z bezpieczeństwem zostało sklasyfikowane, a w razie potrzeby kwalifikowane. Logika jest prosta: kompilator, który po cichu wprowadza błąd, albo narzędzie testujące, które pomija przypadki,

podważa całą argumentację bezpieczeństwa zbudowaną na jego wynikach. Klasyfikacja opiera się na ocenie tego, jak bardzo błąd narzędzia może zaszkodzić i jak prawdopodobne jest, że zostanie wykryty innymi środkami.

Pytanie, na które branża nie ma dziś wspólnej odpowiedzi, brzmi: czy asystent generujący kod jest takim narzędziem. Argument za: wpływa bezpośrednio na treść kodu produkcyjnego. Argument przeciw: nie jest deterministyczny, więc kwalifikacja w klasycznym rozumieniu – czyli wykazanie, że dla danej klasy danych wejściowych daje poprawny wynik – jest metodologicznie niewykonalna.

Praktyka, jaka się wykształca, omija ten problem, a nie rozwiązuje go. Skoro nie da się kwalifikować generatora, kwalifikuje się to, co stoi za nim. Narzędzia analizy statycznej z certyfikatem jednostki notyfikowanej dla norm IEC 61508, ISO 26262, IEC 62304 czy EN 50128 są dostępne od lat i to one przejmują ciężar dowodowy. Kod z modelu traktuje się wówczas dokładnie tak, jak kod od nowego pracownika, którego nikt jeszcze nie zna: wpuszcza się go do procesu na normalnych zasadach i weryfikuje w całości.

Cztery pytania i stan na dziś

W literaturze branżowej powtarza się od dwóch lat zestaw czterech pytań bez rozstrzygnięcia. Warto je przywołać wraz z tym, co zmieniło się w międzyczasie.

Kto odpowiada za błąd w kodzie wygenerowanym przez model? Na gruncie rozporządzenia CRA odpowiedź jest już jednoznaczna i przestała być przedmiotem dyskusji: odpowiada producent wprowadzający wyrób na rynek. Otwarte pozostaje jedynie rozliczenie wewnątrz łańcucha dostaw, na zasadach umownych.

Czy taki kod można certyfikować zgodnie z normami? Tak, na tych samych zasadach co każdy inny, przy czym MISRA C:2025 wprost odbiera mu ulgi przewidziane dla kodu generowanego automatycznie. Certyfikacji podlega proces i wynik, nie zaś pochodzenie kodu.

Czy licencja narzędzia jest zgodna z licencją jego projektu? Pytanie pozostaje otwarte i jest dziś bardziej złożone niż dwa lata temu, ponieważ dochodzi

Kto się podpisuje

Rozporządzenie (UE) 2024/2847 przypisuje obowiązki producentowi wprowadzającemu produkt z elementami cyfrowymi na rynek unijny. Nie dostawcy modelu językowego. Nie autorowi wtyczki do środowiska programistycznego. Nie społeczności projektu otwartego, z którego pochodzi biblioteka.

Producentem jest podmiot, który podpisuje deklarację zgodności – i to on w ciągu doby od potwierdzenia aktywnego wykorzystania podatności musi wskazać, które wyroby i w których wersjach ją zawierają. Dotyczy to również sytuacji, w której podatność pochodzi z komponentu dostarczonego przez kogoś innego.

Zakres terytorialny obejmuje także producentów spoza Unii, jeżeli ich wyroby trafiają na rynek unijny.

Tabela 2. Kalendarz CRA i stan gotowości po stronie firmware'u

Termin	Co zaczyna obowiązywać	Co musi być gotowe w zespole firmware
11.06.2026	Rozdział IV – organy krajowe i jednostki notyfikowane	Wskazana osoba odpowiedzialna za kontakt z organem; ustalona ścieżka zgłoszeniowa
11.09.2026	Zgłaszanie aktywnie wykorzystywanych podatności i poważnych incydentów: 24 h/72 h/14 dni	Wykaz komponentów dla każdej wydanej wersji firmware'u; zestawianie wykazu z bazami podatności; rejestracja na platformie ENISA; polityka skoordynowanego ujawniania podatności; procedura wewnętrzna mieszcząca się w oknie 24 godzin
11.12.2027	Wymagania bezpieczeństwa w projektowaniu, ocena zgodności, dokumentacja techniczna, oznakowanie CE, SBOM, obowiązek aktualizacji przez okres wsparcia	Generowanie wykazu komponentów jako krok potoku budowania; analiza zagrożeń; bezpieczny rozruch i podpisywanie aktualizacji zdalnych; zadeklarowany i obstużony okres wsparcia, co do zasady nie krótszy niż pięć lat; kompletna dokumentacja techniczna

do niego kwestia licencji samych modeli otwartych, integrowanych bezpośrednio w wyrobach – a takich integracji dokonuje, według badania Black Duck, dziewięćdziesiąt sześć procent respondentów.

Czy dostawca modelu może wykorzystać mój kod jako dane uczące? Odpowiedź zależy wyłącznie od warunków konkretnej usługi i od wybranego planu. Dla firm pracujących nad wyrobami objętymi tajemnicą handlową albo ograniczeniami eksportowymi jest to pytanie, które trzeba rozstrzygnąć przed wdrożeniem narzędzia, a nie po nim.

8. Warsztat: jak to robić dobrze

Pora zebrać wszystko w procedurę możliwą do wdrożenia w normalnej firmie, bez zatrudniania osobnego zespołu.

Reguła pierwsza i nadrzędna

Generator nie może być jedynym walidatorem. Zasada ta wygląda na oczywistą, a łamana jest nieustannie, ponieważ ma postać wygodną i niewinną: pytamy model, czy jego własny kod jest poprawny. Model odpowiada, że tak, i odpowiada przekonująco, ponieważ ocenia go tym samym rozumowaniem, którym go wytworzył. Jeżeli w rozumowaniu tkwił błąd, sprawdzenie go nie wykryje – powtórzy.

Niezależna weryfikacja oznacza narzędzie o innej naturze: kompilator, analizator statyczny, test wykonujący kod, pomiar na przyrządzie. Może to być również narzędzie wykorzystujące sztuczną inteligencję, pod warunkiem że nie jest to ten sam model działający na tym samym kontekście.

Specyfikacja zamiast prośby

Różnicę między poleceniem dobrym a złym najlepiej widzieć na przykładzie, który powtarza się w literaturze przedmiotu. Polecenie „napisz program do obsługi czujnika” daje niewiele. Polecenie „wygeneruj funkcję w języku C, odczytującą czujnik temperatury TMP36 przez magistralę I²C i obsługującą przypadki błędne” daje zwykle kod używalny za pierwszym razem.

Dla firmware'u warto pójść dalej i zamiast polecenia sformułować **kontrakt behawioralny**, czyli opis tego, po czym poznamy, że wynik jest dobry, wyrażony

w kategoriach mierzalnych. Nie „ma działać szybko”, tylko „obsługa przerwania nie może przekroczyć dwunastu mikrosekund w najgorszym przypadku”. Nie „ma oszczędzać energii”, tylko „pobór prądu w trybie uśpienia poniżej ośmiu mikroamperów”. Nie „ma poprawnie odczytywać czujnik”, tylko „dla wartości surowej 0xFF80 funkcja podaje minus osiem setnych stopnia”. Nie „ma być odporne na błędy”, tylko „przy zwarcu linii danych do masy funkcja wraca w czasie poniżej pięćdziesięciu milisekund, ze statusem błędu”.

Taki opis ma dwie zalety naraz. Po pierwsze, model dostaje ograniczenia, które rzeczywiście zawężają przestrzeń rozwiązań. Po drugie – i jest to ważniejsze – kryteria akceptacji sformułowane w ten sposób **są już gotowymi przypadkami testowymi**. Praca włożona w specyfikację nie jest pracą dodatkową; jest tą samą pracą, którą i tak trzeba wykonać, tyle że wykonaną wcześniej.

Pętla zamknięta

Właściwy przebieg pracy wygląda tak: specyfikacja, testy, implementacja, budowanie, wgranie na sprzęt, pomiar, decyzja o kolejnym obiegu. Jest to dokładnie ta pętla, którą modelują opisane w rozdziale trzecim testy sprzętowe – z tą różnicą, że tam ocenia ją oprogramowanie badawcze, a u nas oceniacie ją wy.

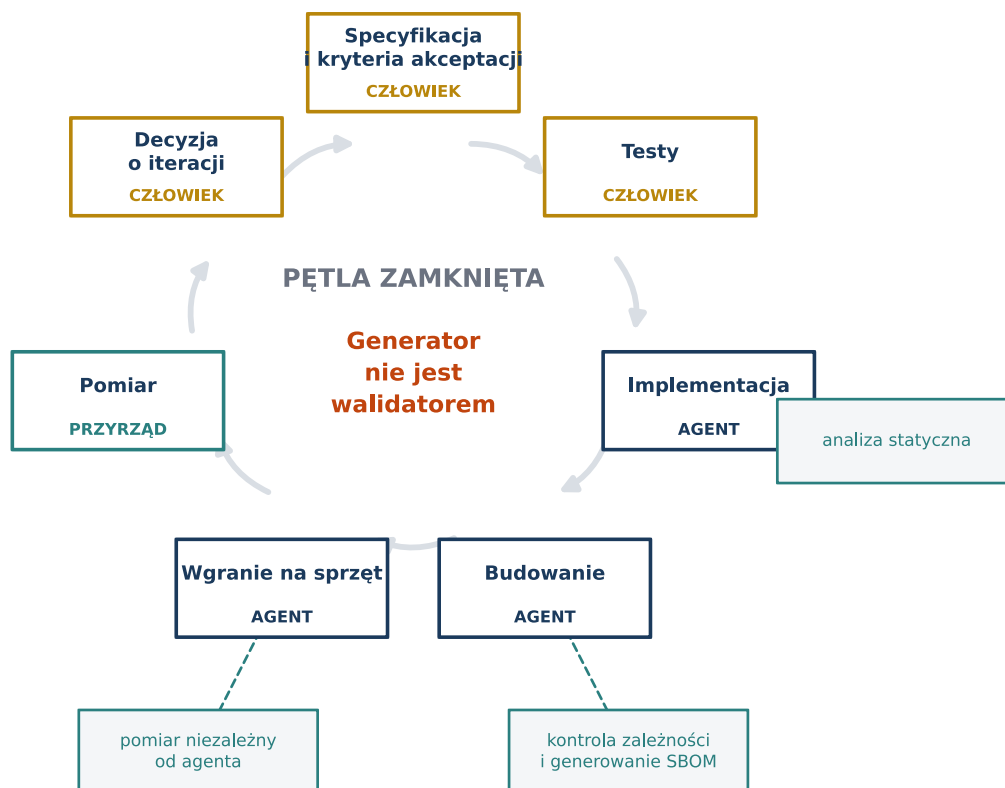
Kluczowa jest kolejność dwóch pierwszych kroków. Testy powstają przed implementacją i nie pisze ich ten sam przebieg modelu, który będzie pisał kod. Jeżeli agent generuje najpierw kod, a potem testy do niego, wytwarza testy potwierdzające to, co zrobił – łącznie z błędami.

Stanowisko sprzętowe

Testowanie ze sprzętem w pętli było dotąd praktyką zarezerwowaną dla większych organizacji, głównie ze względu na koszt zbudowania i utrzymania stanowiska. Zmienia się to z dwóch stron.

Od strony komercyjnej pojawiły się platformy budowane wokół agenta, jak przedstawiona w czerwcu bieżącego roku BootLoop Test. Agent pobiera kontekst z plików projektowych obwodu drukowanego, schematów, łańcucha narzędzi, istniejącego kodu oraz logów i generuje

Specyfikacja, kryteria akceptacji i decyzja pozostają po stronie człowieka; bramki kontrolne działają niezależnie od modelu



Rysunek 9. Pętla zamknięta pracy z agentem nad firmware. Specyfikacja, kryteria akceptacji i decyzja pozostają po stronie człowieka

na tej podstawie testy weryfikujące zachowanie układu na poziomie rejestrów. Platforma obsługuje oscyloskopy, analizatory stanów logicznych, sondy debugowe i mierniki poboru mocy, a producent deklaruje zgodność z amerykańskim reżimem kontroli eksportu ITAR, co otwiera ją na sektory regulowane.

Od strony własnej wystarczy zestaw, którego elementy większość z nas już ma: sonda debugowa obsługiwana przez otwarty serwer kontekstowy oparty na bibliotece probe-rs, zasilacz laboratoryjny z pomiarem prądu,

analizator stanów logicznych, ramy testowe w Pythonie oraz zadanie w potoku ciągłej integracji uruchamiane na fizycznej płytce. Wartość nie leży w wyrafinowaniu narzędzi, tylko w tym, że **ocena przychodzi z pomiaru, a nie z opinii modelu o własnej pracy.**

Przegląd nastawiony na to, czego model nie widzi

Przegląd kodu tworzonego maszynowo wymaga innej uwagi niż przegląd kodu napisanego przez kolegę i wynika

Tabela 3. Co zlecać, a czego nie

Wysoki zysk, ryzyko dające się kontrolować	Koszt weryfikacji przewyższa zysk
Szkielety uruchamiania peryferiów i konfiguracja warstwy sprzętowej	Procedury obsługi przerwań o twardych ograniczeniach czasowych
Parsery protokołów komunikacyjnych o udokumentowanej strukturze	Ścieżki objęte wymaganiami bezpieczeństwa funkcjonalnego
Testy jednostkowe uruchamiane na komputerze, wraz z przypadkami brzegowymi	Implementacje kryptograficzne – zawsze biblioteka sprawdzona, nigdy własna
Dokumentacja, komentarze, opisy interfejsów	Sterowniki wymagające uwzględnienia błędów układu wymienionych w erracie
Wyjaśnianie cudzego kodu i kodu zastanego	Kod podlegający certyfikacji, gdzie każda linia wymaga powiązania z wymaganiem
Przenoszenie kodu między rodzinami mikrokontrolerów jako punkt wyjścia	Optymalizacja pod kątem poboru mocy bez pomiaru na sprzęcie
Skrypty budowania, konfiguracja potoku, narzędzia pomocnicze	Obsługa pamięci nieulotnej i procedury aktualizacji zdalnej
Analiza logów, korelacja zdarzeń, wstępna diagnoza	Wszystko, co dotyczy bezpiecznego rozruchu i kluczy kryptograficznych

Checklista przeglądu firmware'u wygenerowanego przez model

Zgodność ze sprzętem

1. Czy wszystkie nazwy rejestrów i pól istnieją w nocie katalogowej tego układu, a nie układu pokrewnego?
2. Czy adresy urządzeń na magistrali są w formacie, jakiego oczekuje użyty interfejs programowy?
3. Czy stan aktywny każdego wyprowadzenia zgadza się ze schematem ideowym?
4. Czy kolejność uruchamiania uwzględnia włączenie zegara peryferium przed jego konfiguracją?
5. Czy sprawdzono erratę układu pod kątem użytych peryferiów?

Czas i współbieżność

6. Czy każda pętla oczekiwania na warunek sprzętowy ma limit czasu?
7. Czy zmienne dzielone z procedurą przerwania są zadeklarowane jako `volatile`, a dostęp do struktur chroniony?
8. Czy sekwencje odczytu, modyfikacji i zapisu są chronione tam, gdzie rdzeń nie ma instrukcji dostępu wyłącznego LDREX i STREX?
9. Czy oszacowano najgorszy czas wykonania w ścieżkach krytycznych?

Zasoby

10. Czy nie ma alokacji dynamicznej ani dużych buforów na stosie?
11. Czy sprawdzono w raporcie konsolidatora, o ile urosły sekcje?
12. Czy w kodzie nie pojawiła się arytmetyka zmiennoprzecinkowa na rdzeniu bez sprzętowej jednostki zmiennoprzecinkowej?

Błędy i jakość

13. Czy każdy status podawany przez funkcje biblioteki jest sprawdzany?
14. Czy nie ma konstrukcji tłumiących błęd zamiast go obsługiwać?
15. Czy ten kod powiela coś, co już istnieje w projekcie?

Zależności

16. Czy każda nowa zależność została świadomie zatwierdzona i czy naprawdę istnieje?
17. Czy wykaz komponentów został wygenerowany po tej zmianie?

to wprost z danych przytoczonych w rozdziale drugim. Kod z modelu jest równo sformatowany, opatrzony komentarzami i ma sensowne nazwy, a więc spełnia wszystkie powierzchowne kryteria, po których zwykle rozpoznajemy staranność. Sprawia to, że **wygląda na lepszy, niż jest, i przyspiesza przegląd dokładnie tam, gdzie powinien go spowolnić**.

Uwagę warto więc przenieść na kategorie, w których wskaźniki jakości pogorszyły się najbardziej. Czy ten kod korzysta z tego, co już mamy, czy dopisał sobie własną wersję? Czy mieści się w przyjętej architekturze, czy tworzy równoległość? Czy obsługuje przypadki brzegowe, czy je pomija? Czy nie zawiera konstrukcji tłumiących błęd zamiast go obsługi – pustego bloku obsługi, zignorowanego kodu powrotu, pętli bez limitu czasu?

9. Rola konstruktora w 2027 roku

Zbierzmy to, co wynika z przytoczonych danych, w obraz i zawodu za rok.

Kompetencja przesuwana się z pisania kodu w stronę **specyfikowania, ograniczania i weryfikowania**. Nie jest to przesunięcie w stronę wymagań mniejszych, lecz w stronę innych. Dowodem najmocniejszym jest tu wynik testu IoT-SkillsBench: procedury napisane przez eksperta dają lepsze rezultaty niż procedury wygenerowane przez model. Wiedza o tym, jak działa magistrala, gdzie w nocie katalogowej szukać ograniczeń czasowych i po czym poznać, że przebieg na ekranie jest niewłaściwy, nie stała się zbędna. Zmieniła miejsce zastosowania: dawniej wchodziła do kodu, dziś wchodzi do opisu tego, jak kod ma powstać i po czym poznamy, że powstał źle.

Materiały ilustracyjne objęte ochroną – odsyłacze do źródeł

Zgodnie z przyjętą zasadą nie odtwarzamy w druku zrzutów ekranu ani materiałów graficznych objętych prawami osób trzecich. Czytelnikom zainteresowanym wyglądem omawianych narzędzi podajemy odsyłacze do źródeł, pod którymi materiały te są dostępne u ich właścicieli:

- Microchip MPLAB AI Coding Assistant, wygląd wtyczki w środowisku Visual Studio Code – <https://www.microchip.com/en-us/tools-resources/develop/mplab-tools-vs-code/mplab-ai-coding-assistant>
- Microchip MCP Server, model danych i sposób podłączenia – <https://www.microchip.com/en-us/resources/model-context-protocol-server>
- Espressif, serwer kontekstowy narzędzi w ESP-IDF 6.0, opis konfiguracji krok po kroku – <https://developer.espressif.com/blog/2026/04/esp-idf-tools-mcp-server/>
- Espressif, praca nad aplikacją zephyrową z asystentem – <https://developer.espressif.com/blog/2026/06/zephyr-coding-with-ai/>
- Nordic Semiconductor, serwery kontekstowe dla nRF Connect SDK i nRF Cloud – <https://www.nordicsemi.com/Products/Technologies/AI-assisted-development>
- Jacob Beningo, samowalidująca się procedura pracy nad plikiem devicetree w systemie Zephyr – <https://www.beningo.com/zephyr-devicetree-with-ai/>
- METR, komunikat o zmianie metodyki badania wydajności, wraz z pełnym zbiorem danych – <https://metr.org/blog/2026-02-24-uplift-update/>
- Veracode, raport GenAI Code Security 2026 – <https://www.veracode.com/resources/analyst-reports/2026-genai-code-security-report/>
- GitClear, badanie utrzymywaności kodu 2026 – https://www.gitclear.com/the_ai_code_quality_maintainability_gap

Ryzyko leży po drugiej stronie tej samej zmiany. Erozja kompetencji, którą sygnalizują wskaźniki utrzy-
mywalności – spadek przebudowy kodu o siedemdzie-
siąt procent i utrzymania kodu zastanego o siedemdzie-
siąt cztery – dotyka najmocniej tych, którzy wchodzą
do zawodu teraz. Inżynier, który nigdy nie musiał samo-
dzielnie przeczytać noty katalogowej ani znaleźć przy-
czyny zawieszenia na analizatorze stanów logicznych,

nie zweryfikuje agenta. Nie dlatego, że jest gorszy, tylko
dlatego, że weryfikacja wymaga wiedzy, którą nabywa
się wyłącznie przez wykonanie tej pracy ręcznie odpo-
wiednią liczbę razy. Odpowiedzialność za utrzymanie tej
ścieżki spoczywa na zespołach, a nie na osobach począt-
kujących, ponieważ to zespół decyduje, co komu zleca.

Umiejętności, które warto ćwiczyć niezależnie od tego,
jak zmieniają się narzędzia, dają się wymienić krótko.

Słowniczek skrótów użytych w artykule

Skróty rozwinięto w kolejności alfabetycznej. Przy nazwach angielskich podano brzmienie oryginalne, ponieważ w tej postaci występują w dokumentacji i w narzędziach.

AI – Artificial Intelligence, sztuczna inteligencja. W artykule odnosi się do dużych modeli językowych generujących kod, nie do sieci neuronowych uruchamianych na mikrokontrolerze.

ARM Cortex-M – rodzina rdzeni procesorowych firmy Arm przeznaczonych do mikrokontrolerów. Warianty M0 i M0+ są najprostsze i nie mają sprzętowej jednostki zmiennoprzecinkowej; M4 i wyższe zwykle ją mają.

CE – oznakowanie zgodności umieszczane na wyrobie wprowadzanym na rynek Unii Europejskiej. Deklaruje spełnienie wymagań wszystkich przepisów mających do niego zastosowanie.

CRA – Cyber Resilience Act, rozporządzenie (UE) 2024/2847 o horyzontalnych wymaganiach cyberbezpieczeństwa dla produktów z elementami cyfrowymi.

CWE – Common Weakness Enumeration, katalog typów słabości oprogramowania i sprzętu, prowadzony przez korporację MITRE. Każdej pozycji odpowiada numer, na przykład CWE-89 dla wstrzyknięcia zapytania SQL.

CycloneDX – otwarty format zapisu wykazu komponentów oprogramowania, rozwijany w ramach OWASP. Obok SPDX jest jednym z dwóch formatów stosowanych w praktyce do SBOM.

ECCN – Export Control Classification Number, amerykański numer klasyfikacyjny kontroli eksportu, przypisywany wyrobom podlegającym ograniczeniom wywozowym.

ENISA – Agencja Unii Europejskiej ds. Cyberbezpieczeństwa. Prowadzi wspólną platformę, przez którą producenci składają zgłoszenia wymagane rozporządzeniem CRA.

ESP-IDF – Espressif IoT Development Framework, oficjalne środowisko programistyczne dla układów ESP32.

FPU – floating-point unit, sprzętowa jednostka zmiennoprzecinkowa. Jej brak oznacza, że działania na liczbach zmiennoprzecinkowych wykonuje biblioteka programowa, kosztem czasu i rozmiaru kodu.

HAL – hardware abstraction layer, warstwa abstrakcji sprzętu. Biblioteka producenta układu ukrywająca bezpośredni dostęp do rejestrów za wywołaniami funkcji.

HIL – hardware in the loop, testowanie ze sprzętem w pętli. Kod jest sprawdzany na rzeczywistym układzie, a wynikiem testu jest zmierzone zachowanie, nie wynik symulacji.

HumanEval, SWE-bench – testy porównawcze mierzące zdolność modeli do rozwiązywania zadań programistycznych. Pierwszy dotyczy pojedynczych funkcji, drugi zgłoszeń w repozytoriach projektów otwartych. Oba operują niemal wyłącznie na kodzie w Pythonie.

I²C – Inter-Integrated Circuit, dwuprzewodowa magistrala szeregowo do komunikacji układu głównego z czujnikami i innymi układami peryferyjnymi.

ICSE – International Conference on Software Engineering, największa konferencja naukowa poświęcona inżynierii oprogramowania.

ITAR – International Traffic in Arms Regulations, amerykańskie przepisy kontroli obrotu wyrobami o przeznaczeniu wojskowym. Zgodność z nimi bywa warunkiem dopuszczenia narzędzia do projektów obronnych.

JTAG – interfejs służący do programowania i uruchamiania układów scalonych, powszechnie wykorzystywany do podłączenia sondy debugowej.

J-Link, ST-Link – popularne sondy debugowe: pierwsza produkowana przez firmę Segger, druga przez STMicroelectronics.

MCP – Model Context Protocol, protokół, przez który zewnętrzne narzędzia i źródła danych udostępniają asystentowi zestaw funkcji do wywołania wraz z opisem ich działania.

METR – Model Evaluation and Threat Research, niezależna organizacja badawcza zajmująca się pomiarem rzeczywistego wpływu modeli językowych na pracę ludzi.

MISRA – Motor Industry Software Reliability Association, obecnie MISRA Consortium. Wydawca zbiorów wytycznych ograniczających użycie języków programowania w oprogramowaniu o wysokim poziomie integralności.

MPLAB – środowisko programistyczne firmy Microchip dla jej mikrokontrolerów.

OWASP – Open Worldwide Application Security Project, fundacja publikująca między innymi listę OWASP Top 10, czyli zestawienie najpoważniejszych zagrożeń dla aplikacji.

PyPI – Python Package Index, publiczne repozytorium pakietów języka Python. Odpowiednikiem dla języka JavaScript jest npm.

RAM – random access memory, pamięć operacyjna układu, w której przechowywane są dane robocze programu.

RFC 9116 – dokument opisujący plik security.txt, czyli ustandaryzowany sposób publikowania przez producenta danych kontaktowych do zgłaszania podatności.

RISC-V – otwarta architektura zbioru instrukcji procesora, alternatywna wobec architektur zamkniętych.

RoHS – dyrektywa ograniczająca stosowanie substancji niebezpiecznych w sprzęcie elektrycznym i elektronicznym.

RTT – Real-Time Transfer, opracowany przez firmę Segger mechanizm dwukierunkowej wymiany danych między układem a komputerem przez interfejs debugowy, bez zatrzymywania programu.

SBOM – software bill of materials, wykaz komponentów oprogramowania. Czytelna maszynowo lista wszystkich składników, z których zbudowano dany obraz firmware'u, wraz z ich wersjami.

SDK – software development kit, zestaw bibliotek, narzędzi i przykładów dostarczany przez producenta układu.

SPDX – Software Package Data Exchange, format zapisu wykazu komponentów oprogramowania, znormalizowany jako ISO/IEC 5962.

SQL – Structured Query Language, język zapytań do relacyjnych baz danych. Wstrzyknięcie zapytania SQL polega na podaniu danych, które baza wykona jako polecenie.

USENIX Security – jedna z czterech najważniejszych konferencji naukowych poświęconych bezpieczeństwu systemów komputerowych.

XSS – cross-site scripting, wstrzyknięcie skryptu do strony wyświetlanej użytkownikowi. W katalogu CWE oznaczone numerem 80.

Czytanie dokumentacji technicznej ze zrozumieniem, łącznie z erratą układu. Praca z oscyloskopem i analizatorem stanów logicznych, czyli zdolność zobaczenia, co układ robi naprawdę. Myślenie architektoniczne i projektowanie pod kątem testowalności. Budżetowanie zasobów – czasu, pamięci, energii – jako nawyk, a nie czynność wykonywana pod koniec projektu.

I na koniec wróćmy tam, gdzie zaczęliśmy. Osiemdziesiąt trzy i pół procent zespołów ma kod z modelu w wyrobach będących już na rynku. Pięćdziesiąt sześć procent – tyle wynosi zdawalność testów bezpieczeństwa i tyle wynosiła rok temu. Model nie zdejmuję z nikogo odpowiedzialności; przesuwa ją w stronę

tego, kto podpisuje deklarację zgodności. Od 11 września 2026 roku nie jest to już figura retoryczna, tylko termin wynikający z rozporządzenia.

ZAPOWIEDŹ. Cyber Resilience Act od strony warsztatu – w następnym numerze

Punkt siódmy niniejszego artykułu z konieczności traktuje rozporządzenie skrótowo, ponieważ tematem głównym była tu sztuczna inteligencja. Sam CRA zasługuje jednak na osobny materiał i przygotowujemy go do publikacji w przyszłym wydaniu EP.

WM, JT

Bibliografia

Stan odsyłaczy na dzień oddania numeru do druku. Wszystkie pozycje opublikowano między lipcem 2025 a sierpniem 2026 roku, o ile nie zaznaczono inaczej.

Badania naukowe i testy porównawcze

- [1] Babiuch M., Smutný P., „Benchmarking Large Language Models for Embedded Systems Programming in Microcontroller-Driven IoT Applications”, *Future Internet* 18(2):94, MDPI, 11.02.2026. DOI: 10.3390/fi18020094
- [2] „EmbedAgent: Benchmarking Large Language Models in Embedded System Development”, ICSE 2026; preprint: arXiv:2506.11003
- [3] Li Y. i in., „Skilled AI Agents for Embedded and IoT Systems Development”, *SenSys*, 26, ACM; preprint: arXiv:2603.19583. DOI: 10.1145/3786335.3813205
- [4] Zhang Z. i in., „Embedded Arena: Iterative Optimization via Hardware Feedback”, arXiv:2606.16190, 2026; repozytorium: github.com/ubicomplab/embedded-arena
- [5] Morabito R., Wu G., „When the Code Autopilot Breaks: Why Large Language Models Falter in Embedded Machine Learning”, *Computer (IEEE)* 58(11):32–41, 2025; preprint: arXiv:2509.10946
- [6] „Agentic Agile-V: From Vibe Coding to Verified Engineering in Software and Hardware Development”, arXiv:2605.20456, 2026
- [7] „Security Degradation in Iterative AI Code Generation. A Systematic Analysis of the Paradox”, arXiv:2506.11022, 2025 – przegląd obejmujący prace Pearce’a i Perry’ego, cytowane w rozdziale o bezpieczeństwie
- [8] Spracklen J. i in., „We Have a Package for You! A Comprehensive Analysis of Package Hallucinations by Code Generating LLMs”, *USENIX Security Symposium* 2025; preprint: arXiv:2406.10279
- [9] „Re-evaluating LLM Package Hallucinations on the 2026 Frontier”, arXiv:2605.17062, 2026
- [10] Becker J., Rush N., Barnes E., Rein D., „Measuring the Impact of Early-2025 AI on Experienced Open-Source Developer Productivity”, *METR*, arXiv:2507.09089, 07.2025
- [11] *METR*, „We are Changing our Developer Productivity Experiment Design”, 24.02.2026, metr.org/blog/2026-02-24-uplift-update

Raporty branżowe i badania rynku

- [12] RunSafe Security, „2025 AI in Embedded Systems Report: AI Is Here. Security Isn’t.”, 12.2025 – badanie ponad 200 specjalistów z USA, Wielkiej Brytanii i Niemiec
- [13] Black Duck/Censuswide, „The State of Embedded Software Quality and Safety 2025” – badanie 785 specjalistów oprogramowania wbudowanego
- [14] Veracode, „2026 GenAI Code Security Report”, 28.07.2026, veracode.com/resources/analyst-reports/2026-genai-code-security-report
- [15] Veracode, „Spring 2026 GenAI Code Security Update: Despite Claims, AI Models Are Still Failing Security”, 07.2026, veracode.com/blog/spring-2026-genai-code-security
- [16] GitClear, „The Maintainability Gap: 2026 AI Code Quality Research”, gitclear.com/the_ai_code_quality_maintainability_gap – analiza 623 mln zmian z lat 2023–2026
- [17] GitClear, „AI Coding Tools Attract Top Performers – But Do They Create Them?”, 01.2026
- [18] DORA/Google Cloud, „2025 State of AI-assisted Software Development”, 09.2025
- [19] DORA/Google Cloud, „The ROI of AI-assisted Software Development”, 2026
- [20] Stack Overflow, „Developer Survey 2025” – badanie około 49 000 programistów
- [21] Cloud Security Alliance, „AI-Generated Code Vulnerability Surge 2026”, nota badawcza, 04.04.2026
- [22] Cloud Security Alliance, „Slopsquatting: AI Code Hallucinations Fuel Supply Chain Attacks”, 19.04.2026
- [23] Verizon, „Data Breach Investigations Report 2026”

Przepisy, normy i wytyczne

- [24] Rozporządzenie Parlamentu Europejskiego i Rady (UE) 2024/2847 z 23 października 2024 r. w sprawie horyzontalnych wymagań cyberbezpieczeństwa w odniesieniu do produktów z elementami cyfrowymi (Cyber Resilience Act), Dz.U. UE L z 20.11.2024
- [25] Komisja Europejska, wytyczne dotyczące stosowania rozporządzenia (UE) 2024/2847, 27.07.2026
- [26] ENISA, instrukcje rejestracji na wspólnej platformie zgłoszeniowej, 31.07.2026
- [27] MISRA C:2025, „Guidelines for the use of the C language in critical systems”, MISRA Consortium, 03.2025
- [28] MISRA C:2025 Addendum 5, „Coverage of MISRA C:2025 against the Common Weakness Enumeration”, 03.2025, misra.org.uk/app/uploads/2025/03/MISRA-C-2025-ADD5.pdf
- [29] MISRA C:2025 Addendum 6, „Applicability of MISRA C:2025 to the Rust Programming Language”, 03.2025, misra.org.uk/app/uploads/2025/03/MISRA-C-2025-ADD6.pdf
- [30] ISO 26262:2018, „Road vehicles – Functional safety”, w szczególności część 6 (rozwoj oprogramowania) i część 8 (kwalifikacja narzędzi)
- [31] CISA i partnerzy międzynarodowi, „2026 Minimum Elements for a Software Bill of Materials”, 29.07.2026

Narzędzia i materiały producentów

- [32] Microchip Technology, „MPLAB AI Coding Assistant”, microchip.com/en-us/tools-resources/develop/mplab-tools-vs-code/mplab-ai-coding-assistant
- [33] Microchip Technology, „Model Context Protocol Server”, 11.2025, microchip.com/en-us/resources/model-context-protocol-server
- [34] Espressif, „ESP-IDF Tools Local MCP Server: Build, Flash, and Manage Projects from Your AI Assistant”, 28.04.2026, developer.espressif.com/blog/2026/04/esp-idf-tools-mcp-server
- [35] Espressif, „Developing a Zephyr IoT app with AI”, 15.06.2026, developer.espressif.com/blog/2026/06/zephyr-coding-with-ai
- [36] Nordic Semiconductor, „AI-assisted development”, 05.2026, nordicsemi.com/Products/Technologies/AI-assisted-development
- [37] Beningo J., „Zephyr Devicetree with AI: A Self-Validating Skill Guide”, 16.07.2026, beningo.com/zephyr-devicetree-with-ai
- [38] *embedded-debugger-mcp* – serwer MCP dla sond obsługiwanych przez probe-rs, licencja MIT, github.com/adancurusul/embedded-debugger-mcp
- [39] BootLoop, „BootLoop Test” – platforma testowa ze sprzętem w pętli sterowana agentem, premiera 01.06.2026
- [40] Parasoft, „MISRA C 2025: Safer Embedded Code for AI and Rust”, 08.02.2026, parasoftware.com/blog/misra-c-2025-rust-challenges

Stabilizacja pętli trójfazowego prostownika Vienna, część 2

W pierwszej części tej serii pokazałem, jak złożyć dwa modele przełącznika PWM do przewidywania odpowiedzi AC jednofazowego przetwornika T-type. Następnie rozszerzyłem analizę na trójfazowy prostownik Vienna sterowany algorytmem dq0. Po kilku sekwencjach pomiaru AC, trzy pętle zostały ustabilizowane za pomocą klasycznych analogowych kompensatorów typu 2. Symulacje przejściowe potwierdziły stabilne prądy wejściowe o niskiej zawartości harmonicznym zarówno przy niskim, jak i wysokim napięciu sieciowym.

W tej drugiej części będziemy kontynuować symulację prostownika Vienna z modelem cykl-po-cyklu (przełączającym) w celu sprawdzenia całkowitego zniekształcenia harmonicznego (THD) oraz stabilności przy różnych warunkach obciążenia. Następnie pokażę, jak zamknąć pętle za pomocą cyfrowego bloku kompensacji proporcjonalno-całkującej (PI).

Ten ostatni jest cyfrowym odpowiednikiem analogowego kompensatora 2a, który jest podobny do typu 2, lecz nie posiada bieguna wysokoczęstotliwościowego. Jak zobaczymy, dodanie przytrzymania rzędu zerowego (ZOH) do kompensatora cyfrowego osiąga ten sam efekt co brakujący biegun. Omówione zostanie też budowanie kompensatora PI w LTspice przy użyciu wstecznej transformacji Eulera oraz automatyczne generowanie współczynników filtru za pomocą makra. Na koniec symulacje PLECS potwierdzą wyniki.

Symulacje cykl-po-cyklu prostownika Vienna

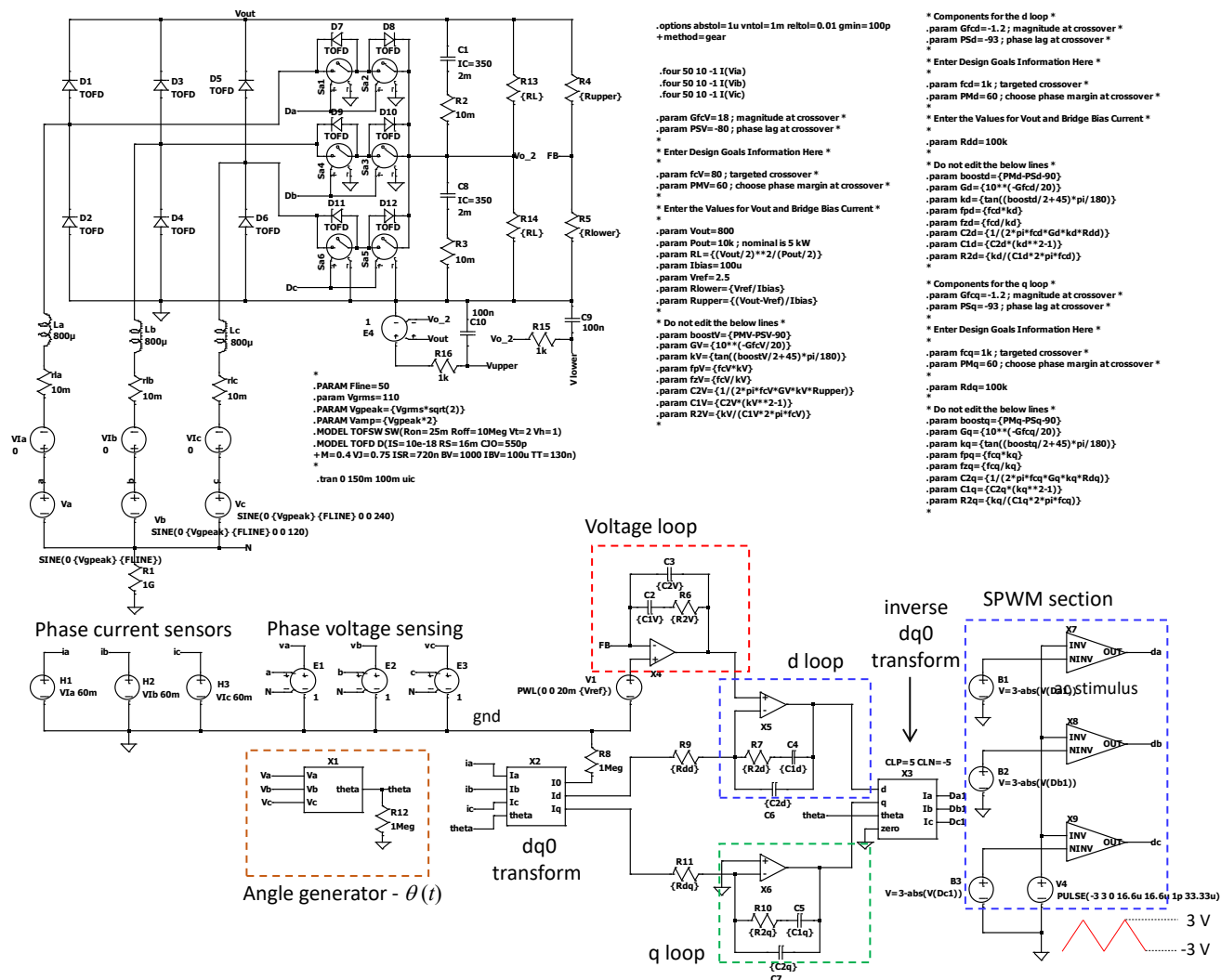
Modele uśrednione są niezerowne pod względem szybkości symulacji, ponieważ nie mają wewnętrznych elementów przełączających. Nie tylko dobrze nadają się do przeprowadzania analiz AC, ale dzięki swoim możliwościom wielkosygnałowym można badać też odpowiedzi przejściowe o dużych sygnałach.

Jednak gdy chodzi o ocenę strat przełączania, naprężeń prądowych i napięciowych lub sprawności, konieczne jest sięgnięcie po wolniejszy model cykl-po-cyklu. Zazwyczaj zaczynam od prostych przełączników sterowanych napięciem dla tranzystorów mocy, a gdy wszystko działa zgodnie z oczekiwaniami, dodaję bardziej kompleksowe modele tranzystorów lub diod.

Kompletny układ widoczny jest na **rysunku 1**. Zastosowałem dwa idealne przełączniki (S_{a1} do S_{a6}) uzupełnione diodą (D_7 do D_{12}). Modulator jest zbudowany z podukładów X_7 do X_9 , które porównują nastawy prądu zrekonstruowane przez blok odwrotnej transformacji dq0 z piłokształtnym napięciem 30 kHz o amplitudzie 6 V (–3 V/3 V).



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.



Rysunek 1. Czas symulacji cykl-po-cyklu dla tego trójfazowego prostownika z idealnymi przełącznikami wynosi około 80 s, co jest dobrym wynikiem dla tak złożonego układu korekcji współczynnika mocy

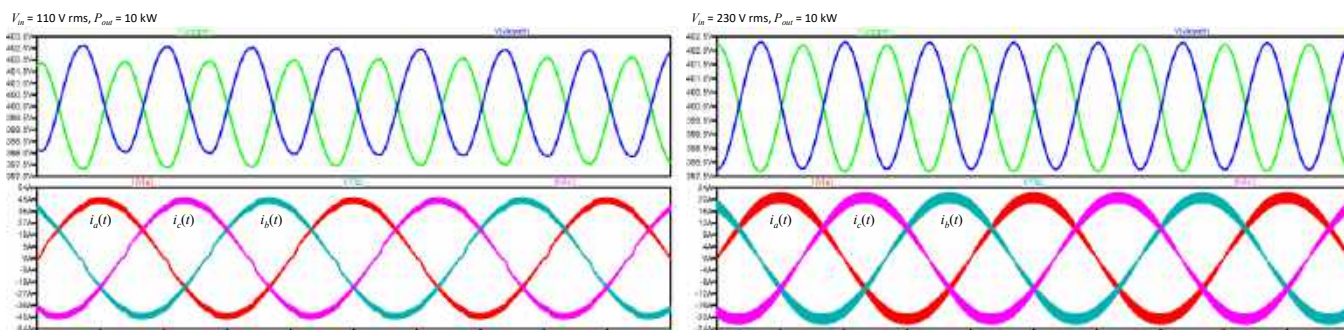
Wyniki symulacji widoczne są na rysunku 2 dla pracy przy obu ekstremalnych napięciach sieciowych z mocą wyjściową 10 kW. Obie szyny są dobrze wyważone na poziomie 400 V, a prąd pobierany z sieci jest sinusoidalny. THD wynosi około 3% w pracy przy niskim napięciu i pozostaje poniżej 6% przy napięciu 230 V rms (400 V linia-do-linii).

Jednak sama dostępność dobrego modelu nie wystarczy do zapewnienia dokładnych wyników symulacji – sprawność przełączania w dużym stopniu zależy od sieci dystrybucji

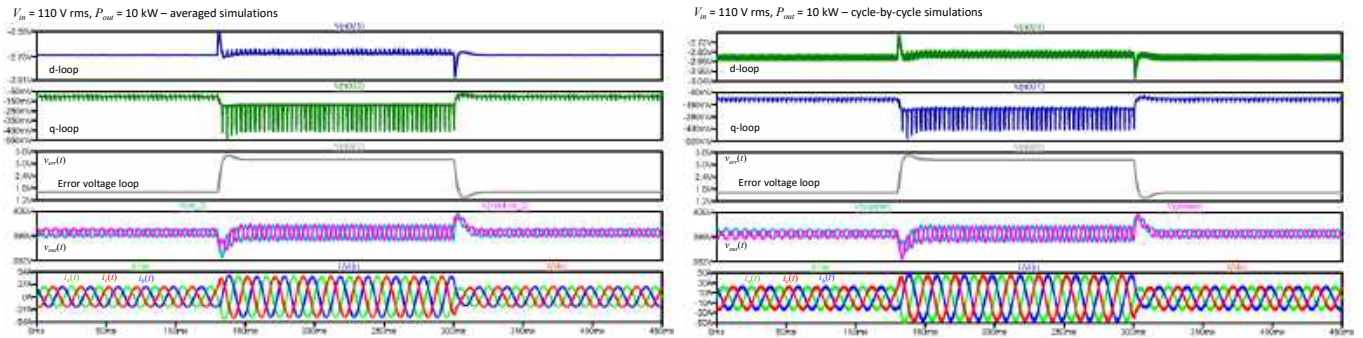
mocy (PDN) i sposobu rozmieszczenia pasywnych w środowisku okablowania. Nieuwzględnienie szeregowych rezystancji i indukcyjności zastępczych kondensatorów łączy DC może sprawić, że symulacja stanie się bezużyteczna.

Odpowiedź przejściowa prostownika Vienna

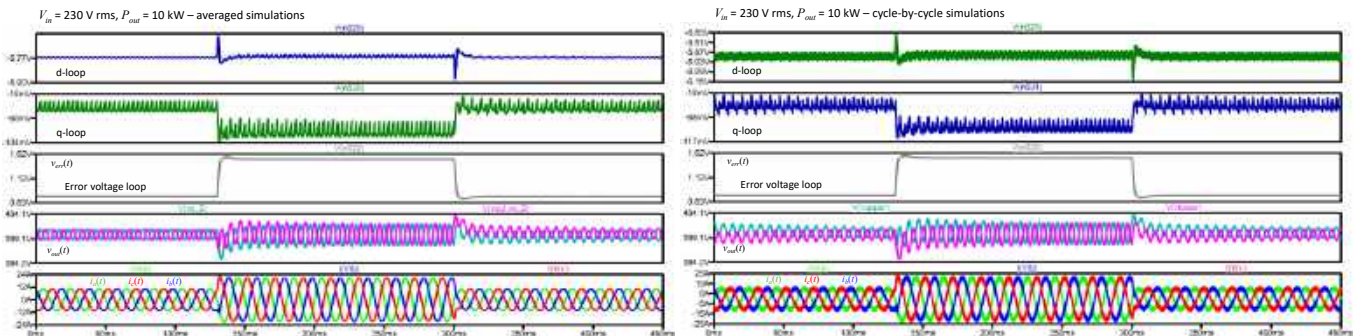
Teraz gdy nasze uśrednione i przełączające układy działają zgodnie z oczekiwaniami, możemy zbadać



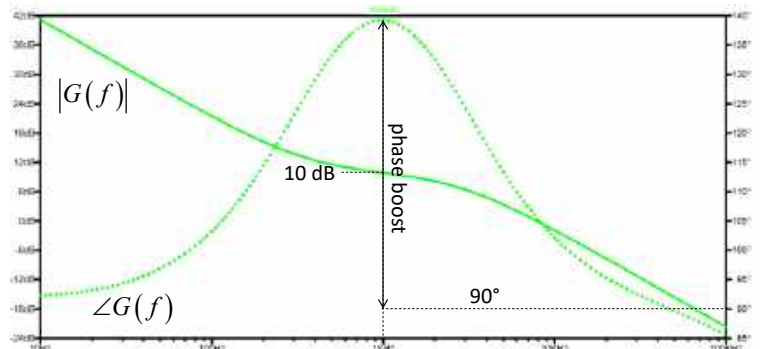
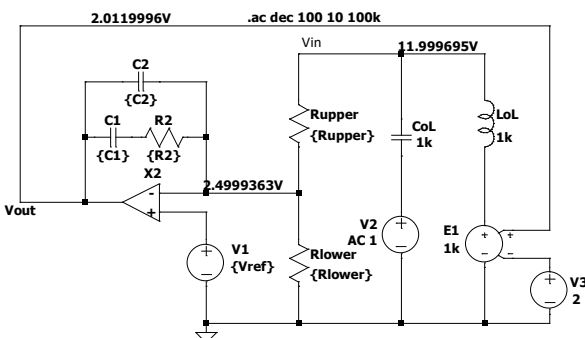
Rysunek 2. Prądy wejściowe w pracy przy niskim napięciu (po lewej) wyglądają dobrze i wykazują THD = 3,2%. W warunkach wysokiego napięcia (po prawej) THD wzrasta do 5,8%



Rysunek 3. Symulacja odpowiedzi przejściowej jest ważnym ćwiczeniem sprawdzającym stabilność pętli. Widoczna jest stabilna praca trzech pętli oraz dobra zgodność przebiegów modeli uśrednionych (po lewej) i przełączających (po prawej)



Rysunek 4. W warunkach skoku obciążenia przy wysokim napięciu sieciowym, odpowiedzi modeli uśrednionych (po lewej) i przełączających (po prawej) są niemal identyczne



Rysunek 5. Kompensator typu 2 zbudowany na wzmacniaczu operacyjnym oferuje potencjalne doładowanie fazowe do 90°

odpowiedź na skok obciążenia. Test odpowiedzi przejściowej stanowi część pakietu analizy stabilności, lecz nie zastępuje pomiaru pętli przeprowadzonego z modelem uśrednionym. Powie nam, czy przyjęta strategia kompensacji utrzymuje dwa napięcia szyny DC 400 V w odpowiednich granicach.

Rysunek 3 zbiera odpowiedzi przejściowe prostownika Vienna pracującego przy niskim napięciu sieciowym. Moc wyjściowa jest schodkowana od 5 do 10 kW w ciągu kilku mikrosekund. Wykresy z lewej strony odpowiadają uśrednionemu układowi implementującemu sześć modeli przełącznika PWM.

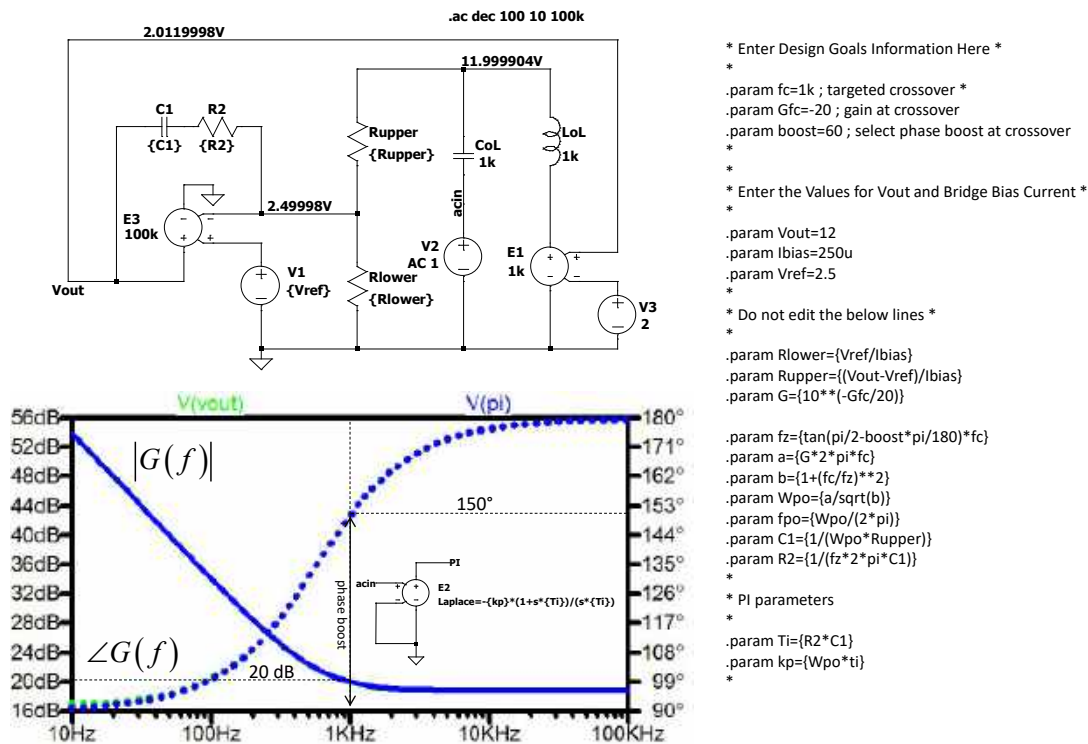
Jak widać z przebiegów, odpowiedź napięciowa jest bardzo płynna, bez znaczącego przeregulowania. Wyjście kompensatora pętli d wykazuje szybszą odpowiedź przejściową, gdyż jego częstotliwość przecięcia wynosi 1 kHz – znacznie powyżej pętli napięciowej. Pętla q nie

zmienia się znacząco, gdzie moc bierna jest celowo utrzymywana na poziomie zera. Rysunek 4 potwierdza doskonałe dopasowanie między uśrednionymi a przełączającymi modelami przy wysokim napięciu sieciowym.

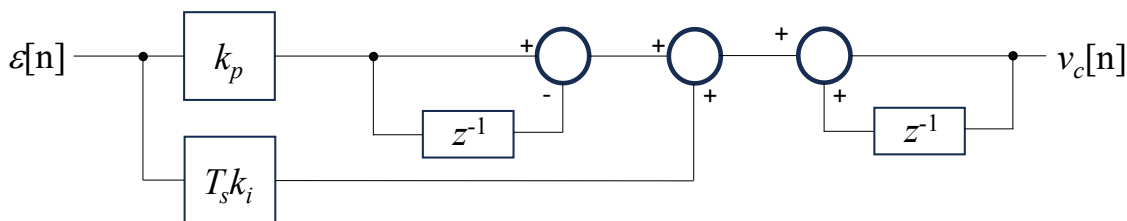
Stabilizacja pętli – analogowe vs. cyfrowe

W części 1 trzy pętle prostownika Vienna zostały ustabilizowane analogowymi kompensatorami typu 2. Ta forma aktywnego filtra zawiera biegun w początku układu i parę biegun-zero. Biegun w początku zapewnia wzmocnienie bliskie nieskończoności przy DC – ograniczone przez wzmocnienie otwartej pętli wzmacniacza operacyjnego A_{OL} – podczas gdy odległość między zerem a biegunem ustala wymaganą ilość doładowania fazowego.

W systemie sterowanym cyfrowo kompensator proporcjonalno-całkujący (PI) jest zwykle stosowany do stabilizacji pętli dq0. Transmitancja takiego filtra wynosi:



Rysunek 6. Kompensator typu 2a nie posiada bieguna wysokoczęstotliwościowego w odróżnieniu od swego kuzyna – typu 2

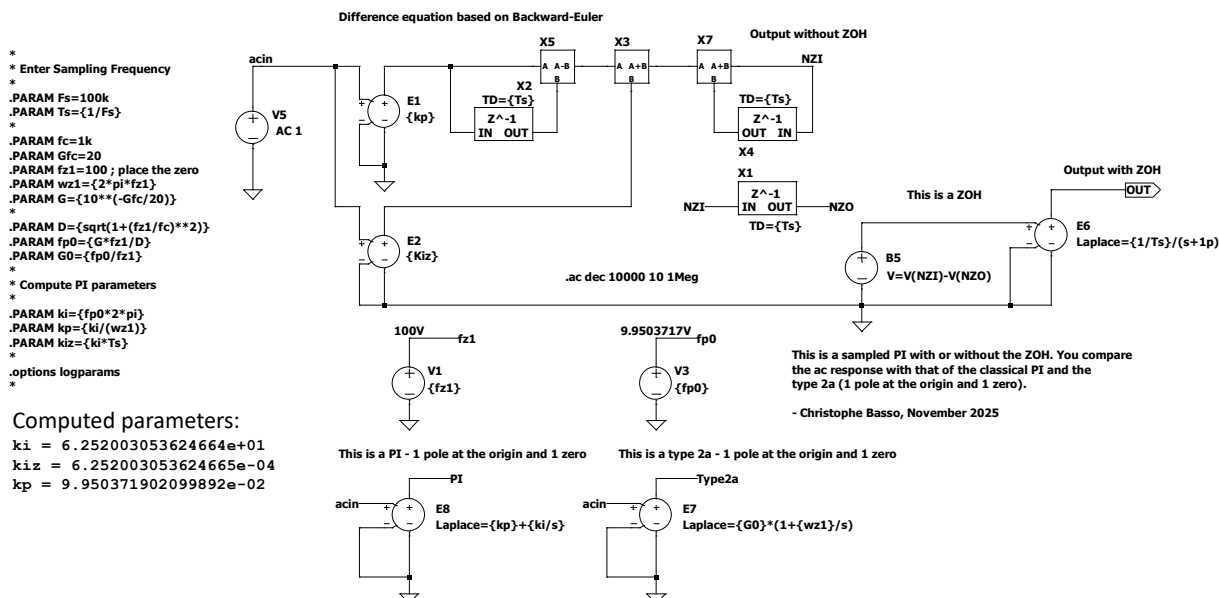


Rysunek 7. Wsteczna transformacja Eulera jest stosowana do transmitancji PI w dziedzinie Laplace'a w celu uzyskania równania różnicowego $v_c[n]$

$$G(s) = k_p \cdot \frac{1+s\tau\tau_i}{s\tau\tau_i}$$

gdzie: $G_0 = \omega_{p0}/\omega_z$, $\omega_z = 1/\tau_i$, $\omega_{p0} = k_p/\tau_i$. Transmitancja ta jest bardzo zbliżona do opisującej kompensator typu 2,

z wyjątkiem braku bieguna wysokoczęstotliwościowego. W istocie opisuje ona tzw. kompensator 2a, który jest analogowym kompensatorem PI.



Rysunek 8. Uproszczony filtr cyfrowy buduje się, składając bloki zgodnie z równaniem różnicowym

Kompensatory cyfrowe

W trójfazowych aktywnych projektach PFC sekcja sterowania zazwyczaj opiera się na pełnej cyfrowej implementacji. Aby zbudować filtr cyfrowy taki jak PI, zaczynamy od transmitancji w dziedzinie Laplace'a i zastępujemy s jego odpowiednikiem przy użyciu wstecznej transformacji Eulera:

$$s \leftarrow \frac{1-z^{-1}}{T_s}$$

Następnie uzyskujemy wyrażenie w dziedzinie z, które konwertujemy na równanie różnicowe opisujące przetwarzanie sygnału $\varepsilon[n]$ w spróbkowaną odpowiedź $v_c[n]$:

$$v_c[n] = k_p \cdot \varepsilon[n] + T_s \cdot k_i \cdot \varepsilon[n] - k_p \cdot \varepsilon[n-1] + v_c[n-1]$$

Po uzyskaniu końcowego wyrażenia w dziedzinie z, modelowanie symulatorem SPICE można realizować na różne sposoby. Jednym z podejść jest opisanie funkcji z^{-1} za pomocą linii opóźniającej. Trzy odpowiedzi częstotliwościowe

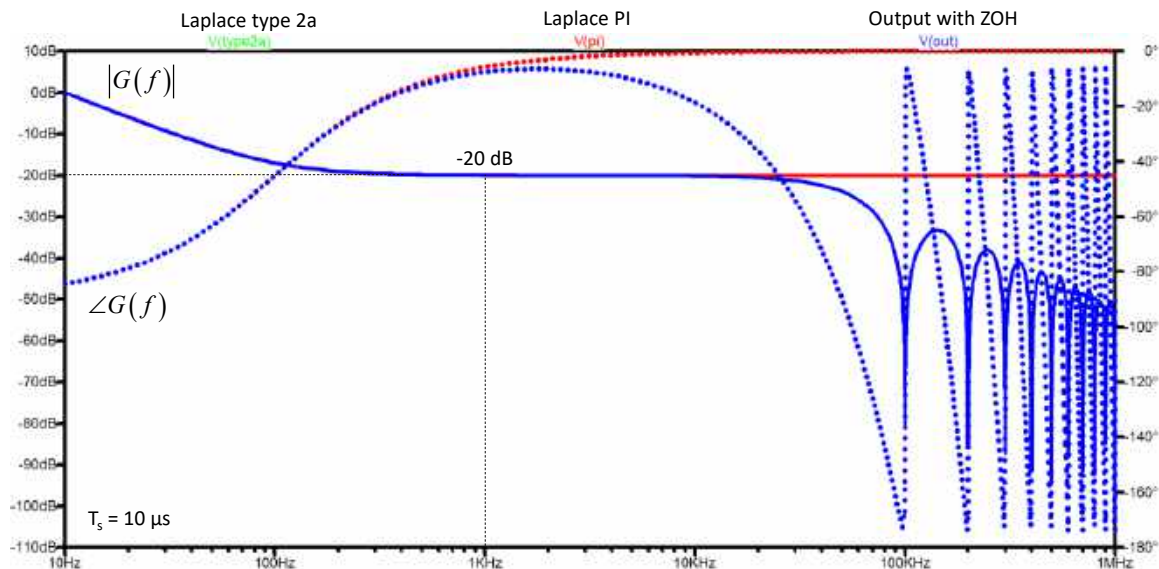
są identyczne pod względem amplitudy do 10 kHz, podczas gdy faza cyfrowa zaczyna się rozbiegać przy około 1 kHz – są to efekty ZOH wprowadzającego opóźnienie $T_s/2$ i kształtową odpowiedź $\sin(x)/x$.

Implementacja w prostowniku Vienna

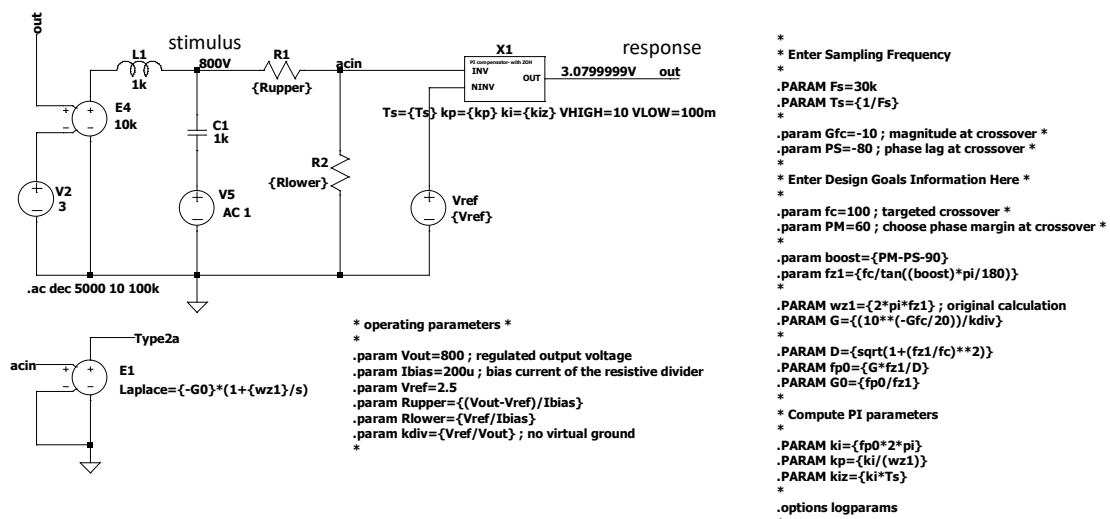
Rysunek 12 pokazuje cały układ wyposażony teraz w filtry cyfrowe. Pętla napięciowa jest kompensowana dla częstotliwości przecięcia 80 Hz, podczas gdy pętla d i q są stabilizowane dla przecięcia 1 kHz. Makro oblicza następujące współczynniki: $k_{pv} = 30,73$; $k_{izv} = 0,297$ dla pętli napięciowej; $k_{pd} = 1,235$; $k_{izd} = 0,143$ dla pętli d; $k_{pq} = 1,039$; $k_{izq} = 0,121$ dla pętli q. Współczynniki te można teraz wprowadzić do graficznego interfejsu użytkownika (GUI) cyfrowego kompensatora.

Symulacje w programie PLECS

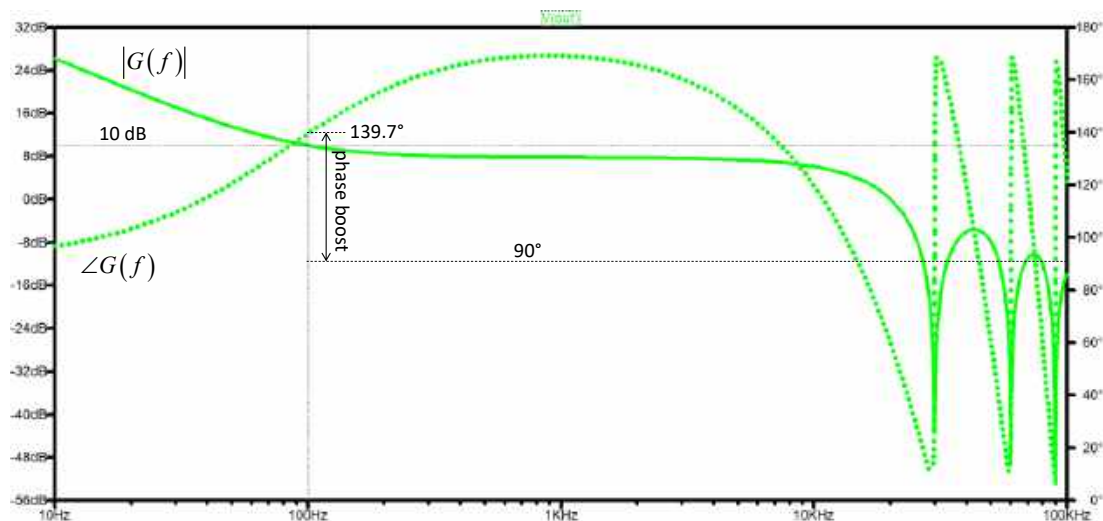
Możliwe jest przeprowadzenie podobnej analizy w programie PLECS – narzędziu do symulacji elektroniki mocy



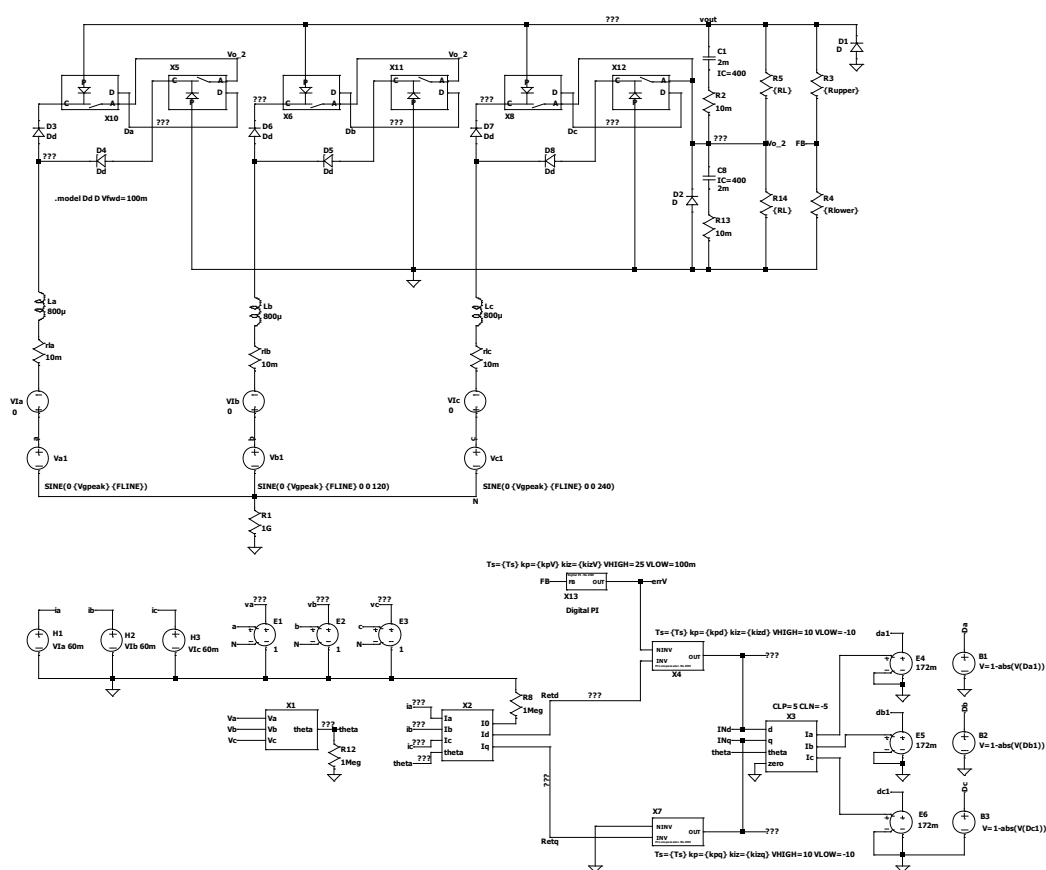
Rysunek 9. Odpowiedź AC potwierdza: wzmacnienie -20 dB przy 1 kHz dla wszystkich trzech wersji filtru



Rysunek 10. Cyfrowy filtr PI jest teraz skojarzony z rezystancyjnym dzielnikiem dla celów regulacji

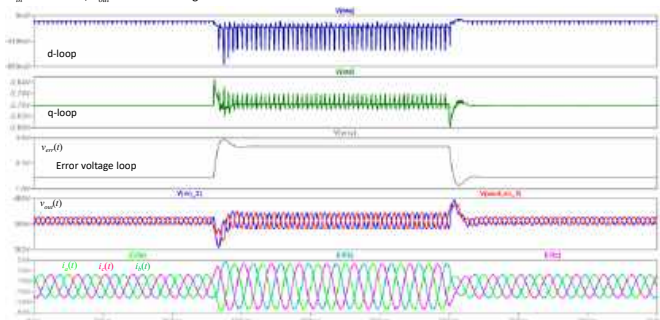


Rysunek 11. Odpowiedź AC odpowiada temu, co wprowadziliśmy do makra: wzmocnienie 10 dB z doładowaniem fazowym 50°

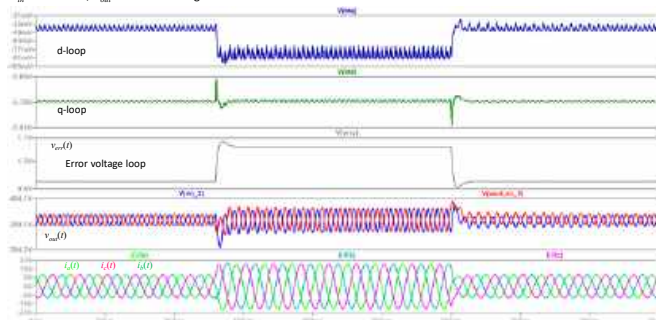


Rysunek 12. Filtry po prostu zastępują wzmacniacz operacyjny i jego elementy pasywne

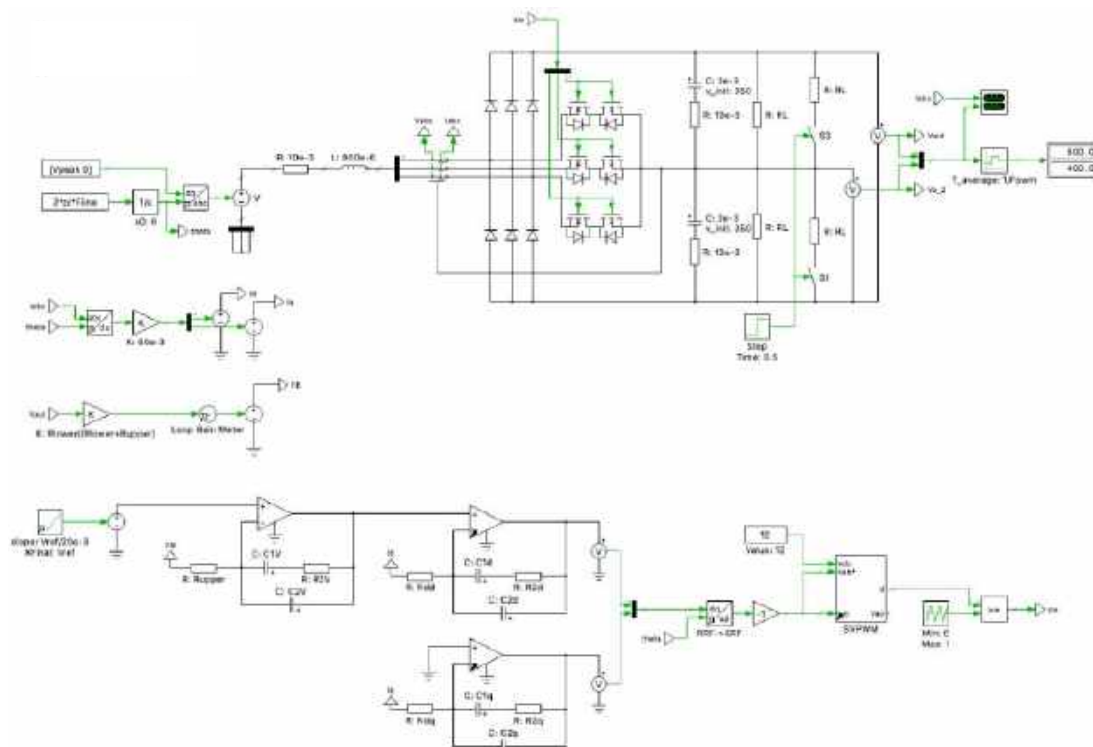
$I_m = 110$ V rms, $P_{avg} = 10$ kW – averaged simulations



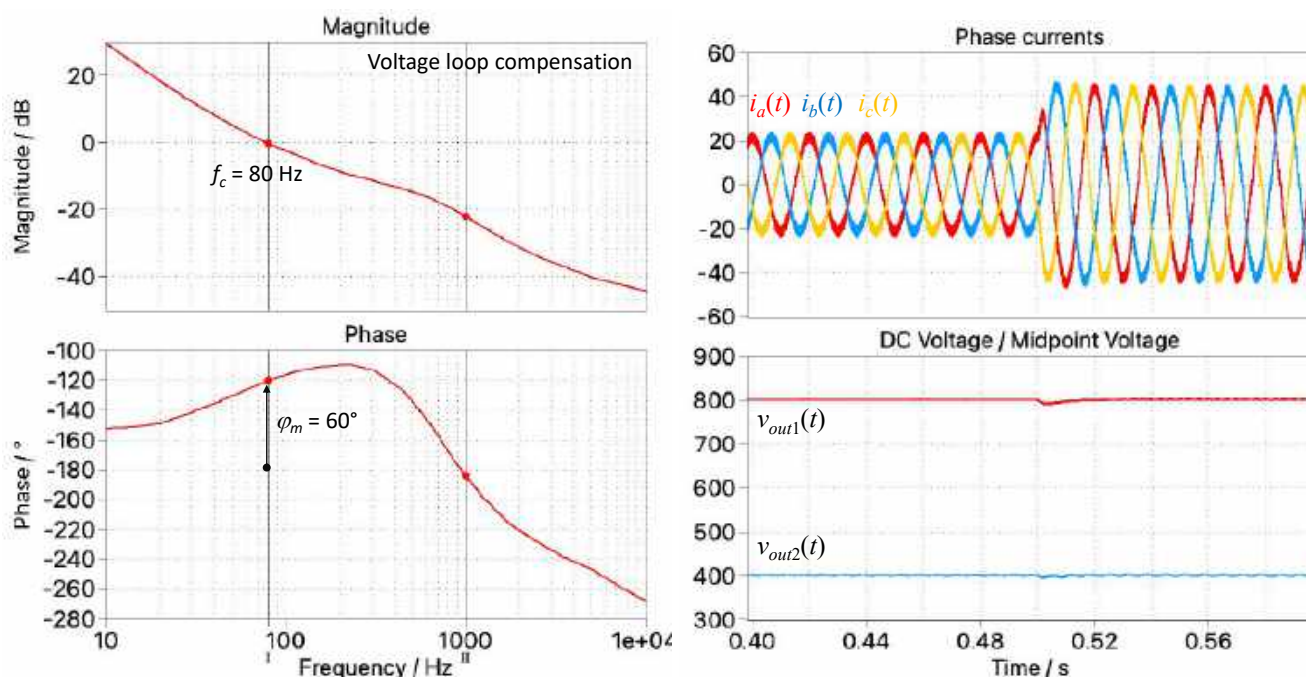
$I_m = 230$ V rms, $P_{avg} = 10$ kW – averaged simulations



Rysunek 13. Odpowiedź przejściowa jest uzyskiwana w czasie poniżej 30 s i potwierdza stabilność przetwornika przy obu poziomach napięcia sieciowego



Rysunek 14. Pierwsza wersja prostownika Vienna w programie PLECS implementuje kompensatory analogowe

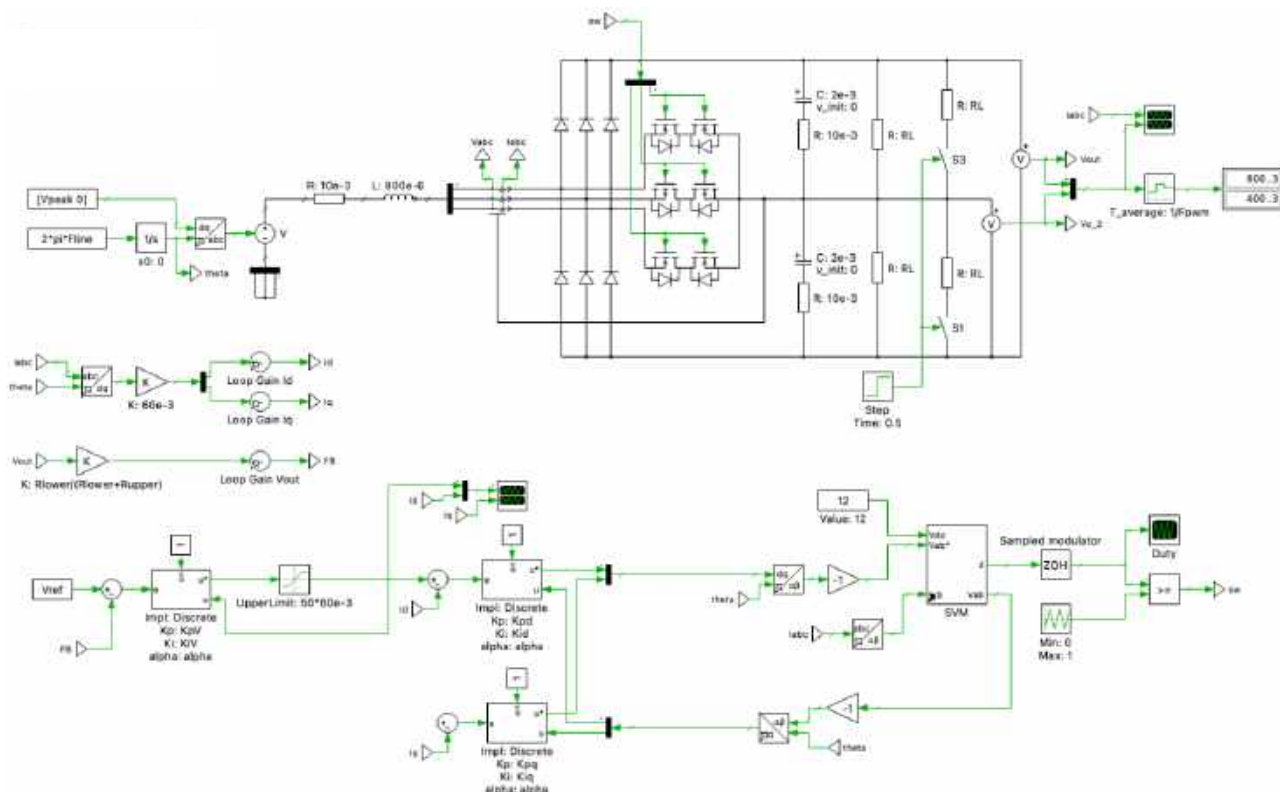


Rysunek 15. Pętla napięciowa jest kompensowana dla częstotliwości przecięcia 80 Hz z zapasem fazy 60°. W przebiegach po prawej stronie odpowiedź przejściowa jest doskonała, z ładnymi sinusoidalnymi prądami wejściowymi

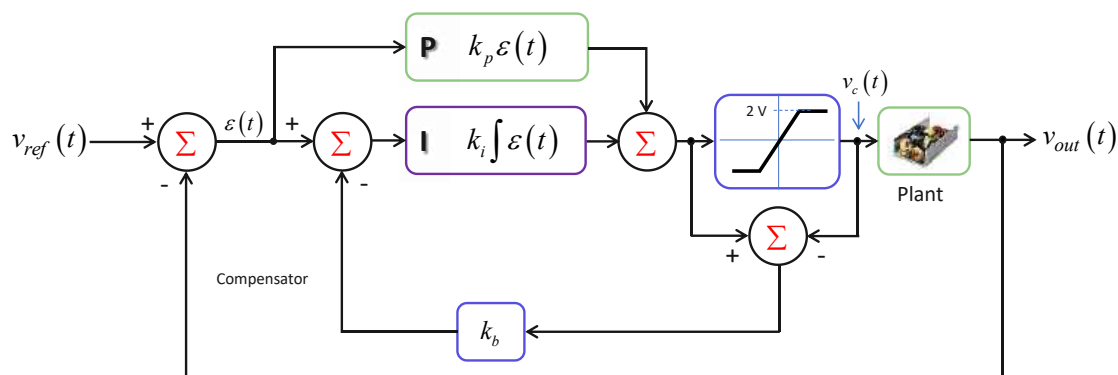
firmy Plexim. Zamiast łączyć elementy pin-po-pinie jak w LTspice, użytkownik umieszcza gotowe bloki opisujące daną funkcję: modulację wektora przestrzennego (SVM), kompensatory PI lub PID w trybie analogowym lub dyskretnym, tranzystory z modelem liniowości odcinkami (PWL) uwzględniającym straty przełączania itd. Beat Arnet, główny inżynier w Plexim U.S., uprzejmie odwzoro- wał moje układy Vienna z kompensatorami analogowymi i spróbkowanymi.

Rysunek 15 pokazuje analizę AC przeprowadzoną przez PLECS dla pętli napięciowej. Kompensacja została usta- wiona na 80 Hz z zapasem fazy 60°. Dwie pozostałe pę- tle, d i q, były stabilizowane dla częstotliwości przecięcia 1 kHz. Odpowiedź przejściowa jest doskonała.

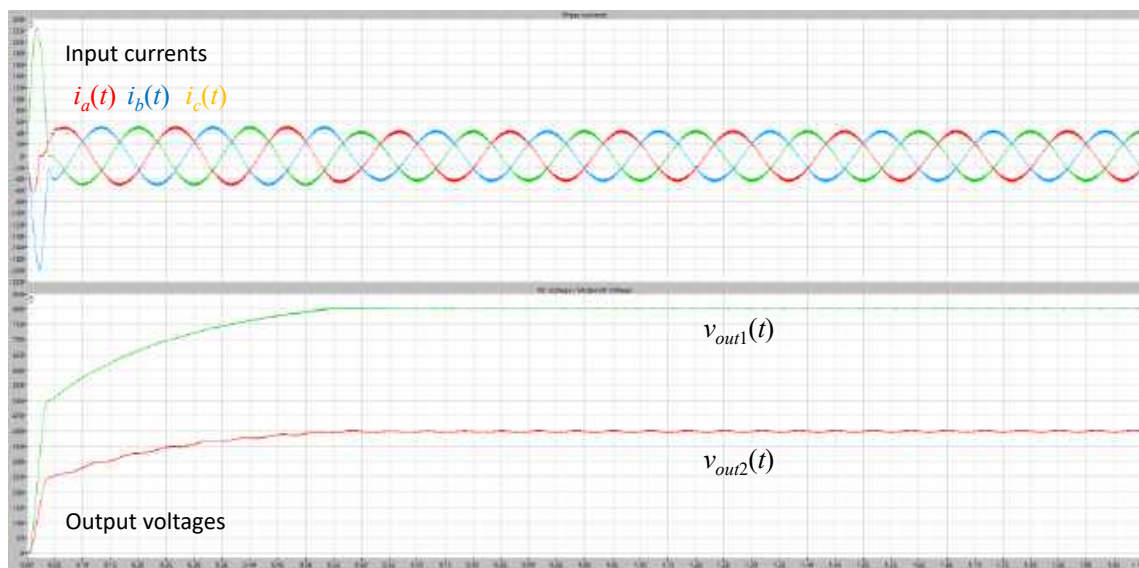
Oryginalny kompensator PI zawiera człon całkujący po- datny na nasycenie podczas rozruchu lub zdarzeń przej- ściowych. Zazwyczaj dodaje się strukturę antywindupową, aby zapobiec tego typu problemom. Napięcie różnicowe jest



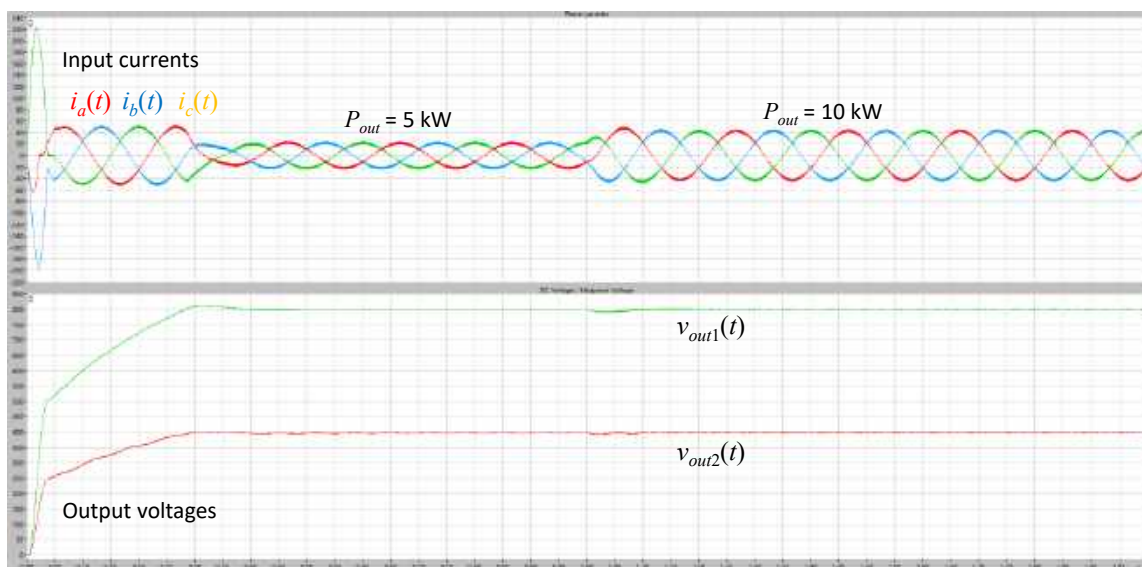
Rysunek 16. Kompensatory w programie PLECS korzystają teraz z próbkowanych wersji kompensatora PI lub typu 2a. Do bloku pętli napięciowej dodano schemat antywindupowy



Rysunek 17. Układ antywindupowy uaktywnia się, aby zapobiec nasyceniu członu całkującego podczas dużego zdarzenia przejściowego



Rysunek 18. Sekwencja rozruchu prostownika Vienna w programie PLECS przebiega płynnie, bez przeregulowania



Rysunek 19. Odpowiedź przejściowa cyfrowo skompensowanego prostownika Vienna symulowana w PLECS jest doskonała, z dobrze kontrolowanym niedoregulowaniem napięć szyn

routowane z powrotem do wejścia integratora za pośrednictwem bloku o współczynniku k_b , którego wartość należy dobrać tak, aby nie degradować odpowiedzi przy aktywnym ograniczniku.

Wnioski

Artykuł zamyka dwuczęściową historię o tym, jak modele uśrednione mogą pomóc w wyznaczeniu odpowiedzi AC trójfazowego prostownika Vienna. Cyfrowy PI korzysta z linii opóźniającej dla operatora z^{-1} i szybko dostarcza oczekiwaną odpowiedź. Po dobrym skalibrowania współczynników symulacje przejściowe potwierdzają,

że przetwornik jest stabilny przy obu ekstremalnych napięciach wejściowych.

Aby dalej zwalidować podejście modelowania, wykazałem, że symulacje PLECS z użyciem moich automatycznych makr również dają stabilną odpowiedź zarówno dla układów analogowych, jak i cyfrowych. Ta technika modelowania może być rozszerzona na inne struktury PFC, jak sześciopak, dostępny jako gotowy szablon w mojej kolekcji plików z lata 2025.

Christophe Basso, Future Electronics, Tuluza,
Francja How2Power Today, kwiecień 2026

Literatura:

1. „Stabilizing The Loops Of A Three-Phase Vienna Rectifier – Part 1” – C. Basso, How2Power Today, marzec 2026.
2. „Double-Pulse Test Helps Validating SiC SPICE Models” – C. Basso, PCIM News Platform, luty 2025.
3. Designing Control Loops for Linear and Switching Power Supplies – C. Basso, Artech House, 2012.
4. „Generic Average Modeling and Simulation of Discrete Controllers” – S. Ben-Yaakov, D. Adar, APEC, Anaheim, CA, 2001.
5. „PID Control of a Vienna Rectifier-Based Power Factor Corrector” – Mathworks, film szkoleniowy, marzec 2021.
6. „PWM Modulator Delay: Analysis verified by PLECS” – B. Arnet, baza wiedzy Plexim, marzec 2026.
7. „Three-Phase 30-kW Vienna PFC Reference Design” – Microchip, 2020.
8. Summer 2025 LTspice Files Collections – ponad 250 szablonów symulacji od autora.

REKLAMA

Mnóstwo doskonałych artykułów: kursów, projektów i nie tylko, przeczytasz na:

EP.com.pl

Ochrona przepięciowa układów (1)

Konstruktor dobiera diodę TVS o napięciu przebicia kilka woltów powyżej napięcia zasilania i uznaje temat ochrony za zamknięty. W karcie katalogowej chronionego układu scalonego widnieje odporność 2 kV w modelu HBM (Human Body Model), norma odbiorcza wymaga 8 kV według IEC 61000-4-2, a instalacja, w której urządzenie ma pracować, przynosi udar 8/20 μ s. Te trzy liczby opisują trzy różne zjawiska, różniące się energią o trzy rzędy wielkości i żadna z nich nie wynika z pozostałych. W cyklu trzech artykułów, publikowanych w trzech kolejnych wydaniach EP, zajmujemy się ochroną przepięciową we wszystkich aspektach istotnych dla praktycznych działań konstruktora. Inspirację i wiedzę czerpiemy ze źródeł sygnowanych przez uznanych w świecie ekspertów w tej dziedzinie. Pierwszy artykuł tego cyklu nie dotyczy elementów ochronnych – dotyczy tego, przed czym mają one chronić. Dopóki zagrożenie nie jest opisane liczbowo, dobór elementu pozostaje zgadywaniem. Elementami ochronnymi zajmujemy się w dwóch kolejnych artykułach tego cyklu.

Przypadek, który nie powinien się zdarzyć

Ten zestaw objawów powtarza się w zgłoszeniach serwisowych na tyle często, że wart jest opisanie jako typowy. Przetwornik pomiarowy z wyjściem prądowym 4...20 mA, pętla zasilana z zasilacza w szafie sterowniczej, na wejściu obwodu pętli prawidłowo dobrana dioda TVS. Wyrób przeszedł badanie odporności na udar według IEC 61000-4-5 na poziomie 3, czyli 2 kV w układzie przewód–ziemia, z wynikiem pozytywnym i z protokołem w dokumentacji. Po instalacji na obiekcie – przepompownia, kilkadziesiąt metrów kabla w ziemi między szafą a przetwornikiem – uszkodzenia wejścia pojawiają się w ciągu kilku tygodni, seryjnie, w kilku egzemplarzach naraz.

Pierwsza myśl: badanie było zbyt łagodne albo dioda została źle dobrana. Obie diagnozy bywają fałszywe. Protokół badania i doświadczenie z obiektu opisują dwa różne zdarzenia. Różnią się impedancją źródła, czasem trwania, drogą wejścia do układu i tym, którędy prąd udarowy wraca do punktu odniesienia. Analizę tego przypadku odkładamy do trzeciej części cyklu, gdzie będziemy mieli już komplet narzędzi. Tutaj służy do postawienia pytania, od którego trzeba zacząć każdy projekt ochrony: co dokładnie wchodzi do układu, którędy i z jaką energią.

Przepięcie, proces przejściowy, przetężenie – porządek pojęciowy

W polskiej praktyce inżynierskiej trzy pojęcia bywają używane zamiennie, choć opisują zjawiska

Co to jest TN-C?

TN-C to system zasilania, w którym funkcje ochronne i neutralne są połączone w jednym przewodzie. Skrót pochodzi z języka francuskiego:

- T (Terre) – punkt neutralny transformatora jest bezpośrednio uziemiony.
- N (Neutre) – części przewodzące instalacji u odbiorcy są połączone z uziemionym punktem transformatora.
- C (Combiné) – funkcje robocze (neutralne N) i ochronne (PE) pełni jeden wspólny przewód PEN.

Dlaczego przerwanie PEN w TN-C niszczy sprzęt?

W normalnych warunkach w gniazdku jest 230 V. Gdy wspólny przewód PEN zostanie przerwany, dochodzi do tzw. „pływania punktu neutralnego”. Napięcie w gniazdku może nagle wzrosnąć nawet do blisko 400 V. Taki stan trwa godzinami (aż do naprawy awarii) i masowo pali urządzenia oraz zabezpieczające je warystory.

Co to jest średnie napięcie (ŚN)?

Średnie napięcie (ŚN) to przedział napięć elektrycznych wyższy niż napięcie w domowych gniazdkach (niskie napięcie), ale niższe niż w liniach dalekobieżnych (wysokie napięcie).

- Zakres: według polskich norm obejmuje napięcia powyżej 1 kV (1000 V) do 30 kV lub 35 kV.
- Zastosowanie: linie napowietrzne i kablowe, które przesyłają prąd z głównych punktów zasilania do transformatorów osiedlowych lub dużych zakładów przemysłowych.

Dlaczego zwarcie na ŚN wywołuje przepięcie w domu?

Transformator osiedlowy obniża napięcie ze średniego (np. 15 000 V) na niskie (400/230 V). Jeśli po stronie średniego napięcia dojdzie do zwarcia z ziemią, potencjał całej ziemi w okolicy transformatora gwałtownie rośnie. To powoduje tzw. przepięcie dorywcze (TOV), które przenosi się do pobliskich domów i niszczy elektronikę.

o różnych czasach trwania i o różnych skutkach dla elementu ochronnego.

- Przepięcie dorywcze (temporary overvoltage, TOV) to podwyższone napięcie o czasie trwania od kilku okresów sieci do sekund, a niekiedy dłużej, nawet do godzin (do momentu naprawy). Typowe przyczyny to przerwa przewodu neutralnego w układzie TN-C, zwarcie doziemne po stronie średniego napięcia albo błąd łączeniowy. Przepięcie dorywcze nie jest procesem przejściowym i to właśnie ono, a nie uderzenie mikrosekundowe, najczęściej niszczy warystory – wrócimy do tego tematu przy okazji ucieczki termicznej.
- Proces przejściowy to zdarzenie jednorazowe o czasie trwania od pojedynczych nanosekund do setek mikrosekund. Do tej kategorii należą wyładowanie elektrostatyczne, paczka szybkich procesów przejściowych, uderzenie łączeniowe i uderzenie piorunowe.
- Przetężenie dotyczy prądu, nie napięcia, i wymaga zupełnie innych środków – bezpiecznika, elementu PTC, ogranicznika elektronicznego. Element ograniczający napięcie nie chroni przed przetężeniem, a element ograniczający prąd nie chroni przed przepięciem. Zestawienie obu funkcji w jednym torze jest tematem trzeciej części.

Drugi podział, równie istotny, dotyczy sposobu, w jaki zaburzenie pojawia się na zaciskach urządzenia. Przepięcie różnicowe (ang. *differential mode*) występuje między przewodami sygnałowymi lub między przewodem fazowym a neutralnym i jest widoczne na wejściu układu wprost. Przepięcie wspólne (ang. *common mode*) występuje między wszystkimi przewodami łącznie a punktem odniesienia – obudową, ziemią, ekranem. Ta druga postać jest w praktyce trudniejsza, bo nie widać jej na oscyloskopie podłączonym w naturalny sposób, to znaczy sondą między przewodami sygnałowymi, a jednocześnie to ona odpowiada za większość uszkodzeń wejść w instalacjach obiektowych.

Systematyczne uporządkowanie tego obszaru dał Ronald B. Standler w monografii z 1989 roku, wznowionej przez wydawnictwo Dover. Książka dzieli materiał na pięć części: objawy i zagrożenia, podstawowe środki zaradcze, typy elementów ochronnych, ich zastosowania oraz weryfikację przyjętych rozwiązań. Przed jej wydaniem literatura przedmiotu była rozproszona po czasopiśmie, patentach, materiałach konferencyjnych i raportach wojskowych. Ten układ pięcioczęściowy jest zresztą powodem, dla którego nasz cykl zaczyna się od zagrożenia, a nie od elementów.

Ochronę dobiera się do statystyki, nie do najgorszego przypadku

Pytanie „jak silne przepięcie może wystąpić” nie ma sensownej odpowiedzi, bo teoretyczna granica leży przy bezpośrednim wyładowaniu w linię zasilającą i żaden element na płycie drukowanej temu nie sprostą. Pytaniem, które

Klasy ekspozycji według IEC 61000-4-5

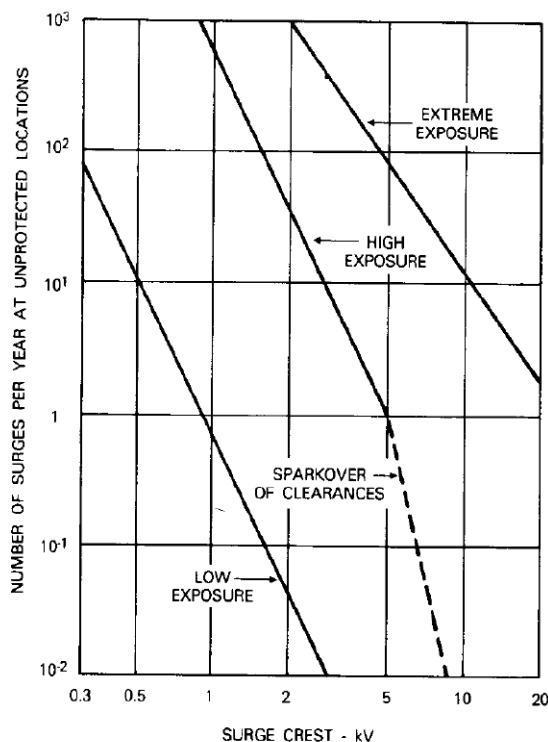
Poziom badania nie jest wybierany dowolnie – norma wiąże go z opisem instalacji, w której urządzenie ma pracować. To jest formalny odpowiednik omawianej wyżej klasyfikacji ekspozycji.

- **Klasa 0** – środowisko dobrze chronione, często wydzielone pomieszczenie: 25 V.
- **Klasa 1** – środowisko częściowo chronione: 500 V.
- **Klasa 2** – instalacja, w której kable zasilające i sygnałowe są rozdzielone, także na krótkich odcinkach: 1 kV.
- **Klasa 3** – instalacja, w której kable zasilające i sygnałowe biegną równolegle: 2 kV.
- **Klasa 4** – instalacja, w której połączenia obejmują kable prowadzone na zewnątrz razem z kablem zasilającym, wykorzystywane zarówno do elektroniki, jak i do obwodów energetycznych: 4 kV.

Klasa 5 obejmuje sprzęt przyłączony do kabli telekomunikacyjnych i napowietrznych linii zasilających na terenie o rozproszonej zabudowie. Klasyfikacja pochodzi z tabeli 1 noty aplikacyjnej STMicroelectronics AN4275, streszczającej wymagania IEC 61000-4-5.

ma odpowiedź, jest: jak często w danym punkcie instalacji występuje przepięcie przekraczające zadany poziom.

François D. Martzloff, który zajmował się tym zagadnieniem kolejno w Southern States Equipment, w General Electric i w NIST, oparł całą metodykę doboru właśnie na takim rozkładzie. Krzywa częstości wystąpień wiąże dwie wielkości: na osi poziomej znajduje się poziom napięcia przepięcia, a na osi pionowej – liczba zdarzeń w ciągu roku, w których przepięcie w danym punkcie instalacji ten poziom przekroczyło. Krzywa opada stromo, co ma bezpośrednią konsekwencję projektową:



Rysunek 1. Liczba zdarzeń w ciągu roku, w których przepięcie w danym punkcie instalacji przekracza wartość odczytaną z osi poziomej. Podstawa: F.D. Martzloff, *Surge Protection Techniques in Low-Voltage AC Power Systems*, NIST

podniesienie odporności urządzenia o jeden stopień eliminuje nieproporcjonalnie dużą część zdarzeń, natomiast dalsze jej podnoszenie daje coraz mniejszy przyrost niezawodności za coraz większe pieniądze.

Do tego dochodzi klasyfikacja ekspozycji. Ten sam wyrób zainstalowany w mieszkaniu w budynku wielorodzinnym w gęstej zabudowie miejskiej i w sterowni na wolno stojącym obiekcie zasilanym linią napowietrzną spotyka się z rozkładami różniącymi się o rząd wielkości. Poziomy badań w normach z rodziny IEC 61000-4 nie są zatem wartościami „dobrymi” ani „złymi” – są odwzorowaniem założonej klasy ekspozycji i tylko w tym kontekście wolno je porównywać.

Siedem zagrożeń w jednej tabeli

Poniższe zestawienie jest podstawowym narzędziem roboczym całego cyklu. Wiersze odpowiadają klasom zagrożeń, które w literaturze i w normach bywają mieszane, choć wymagają odmiennych środków. Wartości energii podano jako oszacowania rzędu wielkości, a nie jako wielkości katalogowe: dla modeli wyładowania elektrostatycznego wynikają z energii zgromadzonej w kondensatorze modelu, czyli z połowy iloczynu pojemności i kwadratu napięcia, natomiast dla udaru 8/20 μ s przyjęto ładunek odpowiadający przybliżeniu trójkątnemu przebiegu i napięcie ograniczania rzędu kilkudziesięciu woltów. Przy innym napięciu ograniczania wynik zmienia

Jak czytać zapis 8/20 μ s

Dwie liczby oddzielone ukośnikiem opisują kształt impulsu w mikrosekundach. Pierwsza to czas czoła, druga – czas liczony od umownego początku impulsu do chwili, w której wielkość opada do połowy wartości szczytowej. Zapis 1,2/50 μ s odnosi się do przebiegu napięcia, zapis 8/20 μ s – do przebiegu prądu.

Dlaczego to są wielkości umowne

Żadnej z nich nie mierzy się wprost. Rzeczywiste czoło impulsu jest zaszumione i oscylacyjne, więc nie da się wskazać na oscylogramie punktu, w którym impuls się zaczyna. Zamiast tego prowadzi się prostą przez dwa punkty odniesienia na zboczu narastającym i ekstrapoluje ją do zera. Dla przebiegu napięciowego czas czoła przyjmuje się jako 1,67 czasu narastania między 30% a 90% wartości szczytowej, dla prądowego – jako 1,25 czasu narastania między 10% a 90%. Metoda pochodzi z normy IEC 60060-1, czyli z klasycznej techniki badań wysokonapięciowych, i została przeniesiona do badań kompatybilności elektromagnetycznej bez zmian.

Co z tego wynika w praktyce

Impuls o „takim samym” oznaczeniu, zmierzony na wyjściu generatora i na zaciskach badanego urządzenia, może mieć zauważalnie różny kształt – i norma to dopuszcza. Na wyjściu generatora dopuszczalne jest ponadto przeregulowanie do 30% wartości szczytowej poniżej zera. Oznaczenie 8/20 μ s opisuje więc klasę przebiegu, a nie pojedynczą krzywą.

się proporcjonalnie. Dwie wielkości z ostatniego wiersza wymagają objaśnienia. LPL I to najwyższy z czterech poziomów ochrony odgromowej zdefiniowanych w normie IEC 62305, czyli założenie najostrzejsze z możliwych. Energia właściwa, wyrażana w megadżulach na om, to całka kwadratu prądu po czasie: mówi, ile ciepła prąd wydzieliby w rezystancji jednego oma, a więc jest bezpośrednią

Tabela 1. Klasy zagrożeń przepięciowych. Zestawienie własne redakcji na podstawie danych liczbowych z przywołanych norm

Zagrożenie	Dokument odniesienia	Przebieg i impedancja źródła	Energia – rząd wielkości	Typowa droga wejścia
ESD komponentowe, model HBM	ANSI/ESDA/JEDEC JS-001	100 pF rozładowane przez 1500 Ω ; narastanie rzędu 10 ns, zanik ok. 150 ns	ok. 0,2 mJ przy 2 kV	wyprowadzenia układu w trakcie montażu i transportu
ESD komponentowe, model CDM	ANSI/ESDA/JEDEC JS-002	rozładowanie ładunku zgromadzonego w samej obudowie; czas poniżej 1 ns, brak rezystora ograniczającego	poniżej 0,1 mJ, ale bardzo duży prąd szczytowy	kontakt wyprowadzenia z powierzchnią o innym potencjale
ESD systemowe	IEC 61000-4-2	150 pF rozładowane przez 330 Ω ; narastanie ok. 0,8 ns; do 8 kV wyładowania kontaktowego	ok. 5 mJ przy 8 kV	złącza i elementy obsługowe dostępne dla użytkownika
Szybkie procesy przejściowe seryjne (EFT/burst)	IEC 61000-4-4	impulsy 5/50 ns w paczkach, częstotliwość powtarzania 5 kHz lub 100 kHz	pojedynczy impuls o małej energii; decyduje liczba powtórzeń	sprężenie pojemnościowe do wiązek kablowych, kleszcze sprzęgające
Udar w instalacji niskiego napięcia	IEC 61000-4-5	1,2/50 μ s przy rozwarciu, 8/20 μ s przy zwarcu; 2 Ω dla portów zasilania, 12 Ω i 42 Ω dla portów sygnałowych	rzędu 0,1–1 J odkładane w elemencie ograniczającym	przewody zasilające i długie linie sygnałowe wewnątrz obiektu
Udar na linii telekomunikacyjnej	ITU-T K.20, K.21; załącznik IEC 61000-4-5	10/700 μ s napięciowo, 5/320 μ s prądowo; impedancja źródła 40 Ω , poziomy do 4 kV	rzędu pojedynczych dżuli	linie wyprowadzone poza obiekt, kable napowietrzne i ziemne
Prąd piorunowy bezpośredni	IEC 62305, EN 61643-11	10/350 μ s; dla najwyższego poziomu ochrony LPL I prąd szczytowy 200 kA, ładunek 100 C	energia właściwa 10 MJ/ Ω dla LPL I	instalacja odgromowa, uziom, przewody wchodzące do obiektu

miarą obciążenia cieplnego elementu, przez który ten prąd przepływa.

Z tabeli 1 wynika obserwacja, którą warto zapamiętać przed lekturą drugiej części: rozpiętość energii między wierszem pierwszym a ostatnim wynosi około dzie więciu rzędów wielkości. Nie istnieje element ochronny, który obsługuje cały ten zakres. Każde rozwiązanie jest kompromisem dobranym do jednego lub dwóch wierszy, a kompletna ochrona urządzenia to zawsze zestaw elementów o różnych właściwościach, ustawionych w określonej kolejności.

Punkt sporny: 8/20 μ s czy 10/350 μ s

Spór o właściwy przebieg do wymiarowania ochronników pierwszego stopnia trwa od lat osiemdziesiątych i ma wymiar wprost finansowy: ogranicznik typu 1 wymiarowany na przebieg 10/350 μ s kosztuje wielokrotność urządzenia wymiarowanego na 8/20 μ s, przy tej samej wartości prądu szczytowego.

Stanowisko, które przeważało w normalizacji europejskiej, wywodzi się z pracy W. Meissena opublikowanej w ETZ w 1983 roku. Autor zaproponował długi przebieg o bardzo dużej zdolności deponowania energii, argumentując, że to on odwzorowuje rzeczywiste obciążenie ogranicznika przy bezpośrednim wyładowaniu; z tej propozycji wyrosła niemiecka norma DIN 0160, a dalej wymaganie prądu udarowego o kształcie 10/350 μ s dla ograniczników typu 1. Linię tę rozwinęła szkoła związana z firmą DEHN, a jej wykładnią po stronie instalacyjnej jest monografia Petera Hassego, przez lata dyrektora zarządzającego tego przedsiębiorstwa i współtwórcy norm ochrony odgromowej.

Stanowisko przeciwne sformułował Martzloff w pracy przedstawionej na konferencji ICLP na Rodos w 2000 roku, dotyczącej rozplywu prądu piorunowego po bezpośrednim wyładowaniu w budynek. Argument jest następujący: prąd wyładowania dzieli się między wszystkie dostępne drogi do ziemi – uziom fundamentowy, instalacje przewodzące, przewody wchodzące do obiektu – więc udział przypadający na pojedynczy ogranicznik jest znacznie mniejszy od wartości całkowitej. Autor zestawiał wyniki modelowania komputerowego z doświadczeniem eksploatacyjnym i na tej podstawie zakwestionował parametry przyjmowane przy wymiarowaniu.

Dla czytelnika projektującego urządzenie, a nie instalację, spór ten ma znaczenie pośrednie, ale realne. To on rozstrzyga, jaka wartość napięcia resztkowego pojawi się na zaciskach zasilania urządzenia po zadziałaniu ogranicznika instalacyjnego – czyli definiuje warunek brzegowy dla ochrony wbudowanej w wyrób. Konstruktor nie ma wpływu na wybór strony sporu, ale ma obowiązek wiedzieć, którego założenia użył przy doborze własnych elementów. Ani jedno, ani drugie stanowisko nie jest błędne; różnią się

odповідzią na pytanie, jaka część prądu wyładowania wchodzi do konkretnej instalacji.

Mit: „dwa kilowolty w karcie katalogowej znaczą, że układ jest odporny”

Zdanie „wejścia zabezpieczone do ± 2 kV HBM” występuje w kartach katalogowych powszechnie i bywa czytane jako deklaracja odporności wyrobu. Nie jest nią. Opisuje warunki, w jakich układ scalony wolno przenosić, montować i lutować, a nie warunki, w jakich gotowe urządzenie ma pracować u użytkownika. Różnicę widać dopiero wtedy, gdy trzy modele wyładowania zestawia się obok siebie.

Model ciała ludzkiego (HBM, opisany w ANSI/ESDA/JEDEC JS-001) odwzorowuje rozładowanie naładowanego człowieka przez rezystancję skóry i dłoni: kondensator 100 pF rozładowany przez rezystor 1500 Ω , czas narastania rzędu 10 ns, zanik około 150 ns. Prąd szczytowy przy 2 kV wynosi zatem około 1,3 A, a energia zgromadzona w kondensatorze modelu – około 0,2 mJ.

Model naładowanego elementu (CDM, opisany w ANSI/ESDA/JEDEC JS-002) opisuje sytuację odwrotną: ładunek gromadzi się w samej obudowie układu, na przykład przy przesuwaniu się w podajniku, i rozładowuje się przez wyprowadzenie stykające się z uziemioną powierzchnią. W obwodzie nie ma rezystora ograniczającego, więc czas narastania schodzi poniżej nanosekundy, a prąd szczytowy przy tym samym napięciu nominalnym jest wielokrotnie większy niż w modelu HBM, przy całkowitym czasie trwania rzędu pojedynczych nanosekund. W nowoczesnym montażu automatycznym to właśnie CDM, a nie HBM, jest dominującym mechanizmem uszkodzeń.

Model systemowy z IEC 61000-4-2 opisuje zupełnie inne zdarzenie: użytkownika dotykającego złącza gotowego wyrobu. Kondensator ma 150 pF, rezystor 330 Ω , czas narastania pierwszego szczytu wynosi około 0,8 ns. Norma określa dla wyładowania kontaktowego 8 kV prąd pierwszego szczytu 30 A, a następnie wartości kontrolne 16 A po 30 ns i 8 A po 60 ns – przebieg ma więc dwie składowe o różnych stałych czasowych. Energia zgromadzona w kondensatorze modelu wynosi tu około 4,8 mJ.

Zestawienie tych trzech opisów daje wynik jednoznaczny: przy przejściu od 2 kV HBM do 8 kV według IEC 61000-4-2 prąd szczytowy rośnie ponad dwudziestokrotnie, energia – ponad dwudziestokrotnie, a czas narastania skraca się o rząd wielkości. Struktura ochronna wbudowana w układ scalony jest wymiarowana na pierwszy z tych przebiegów, bo taki jest cel normy komponentowej. Producent układu nie może zadeklarować odporności systemowej, ponieważ nie zna ani impedancji toru między złączem a wyprowadzeniem, ani geometrii płytki, ani obecności elementów zewnętrznych.

Wyjątkiem są układy, dla których producent wprost podaje odporność wyprowadzeń według IEC 61000-4-2 – wtedy

deklaracja dotyczy już zdarzenia systemowego i wolno się na niej opierać, pamiętając, że dotyczy wyprowadzenia, a nie złącza na obudowie. Granicą między ochroną wewnątrz układu a ochroną na płycie zajmuje się wprost praca zbiorowa pod redakcją Charvaki Duvvury'ego i Haralda Gossnera, poświęcona współprojektowaniu obu poziomów; fizykę samych struktur ochronnych opisuje seria monografii Stevena H. Voldmana, wieloletniego przewodniczącego grupy roboczej ESD Association do spraw pomiarów metodą TLP.

Jak przepięcie wchodzi do układu

Drogi są trzy i warto je rozróżnić, bo każda wymaga innego środka zaradczego. Pierwsza to przewodzenie – zaburzenie wchodzi wprost żyłą zasilającą albo sygnałową i przeciwdziała mu element ochronny umieszczony na tej żyłę. Druga to sprzężenie pojemnościowe, istotne przy zaburzeniach o krótkim czasie narastania, gdzie pojemność rzędu pikofaradów między przewodem zaburzającym a torem sygnałowym wystarcza do przeniesienia impulsu. Trzecia to sprzężenie indukcyjne, czyli napięcie wyindukowane w pętli obwodu przez zmienne pole magnetyczne – i to ona jest najczęściej lekceważona, bo nie widać jej na schemacie.

Rachunek jest prosty i wart przeprowadzenia, bo wynik zaskakuje. Napięcie wyindukowane w pętli równa się iloczynowi jej pola powierzchni i szybkości zmian indukcji magnetycznej, przy założeniu, że pole jest w obrębie pętli w przybliżeniu jednorodne. Przyjmijmy wyładowanie elektrostatyczne o prądzie szczytowym 30 A narastającym w ciągu 1 ns – wartości wprost z normy IEC 61000-4-2 dla 8 kV. W odległości 10 cm od przewodu, którym płynie ten prąd, indukcja osiąga około 60 μ T, a szybkość jej narastania – rząd $6 \cdot 10^4$ T/s. Pętla o polu 10 cm², czyli na przykład ścieżka sygnałowa i jej droga powrotna oddalone o 5 mm na odcinku 20 mm, obejmuje wtedy strumień dający napięcie rzędu kilkudziesięciu woltów.

Kilkadziesiąt woltów w torze zasilanym z 3,3 V, wygenerowane bez żadnego kontaktu galwanicznego, wyłącznie przez geometrię. Wniosek projektowy jest natychmiastowy: dwukrotne zmniejszenie pola pętli zmniejsza wyindukowane napięcie dwukrotnie, a żadna dioda ochronna nie robi tego za nas, bo zaburzenie powstaje już za nią. W wielu układach geometria pętli decyduje o odporności bardziej niż wybór elementu ochronnego – do konsekwencji layoutowych wracamy w części trzeciej.

Osobnym zagadnieniem jest przepięcie w postaci wspólnej. Przepięcie wspólne występuje między wszystkimi przewodami toru łącznie a punktem odniesienia i nie jest widoczne pomiędzy przewodami sygnałowymi. Ma to dwie konsekwencje. Po pierwsze, elementy ochronne włączone różnicowo go nie ograniczają. Po drugie – i to jest źródłem wielu błędnych diagnoz – nie widać go na oscyloskopie podłączonym w sposób naturalny, to znaczy sondą, której masa

Praktyczna konsekwencja

Deklaracja ± 2 kV HBM w karcie katalogowej układu scalonego jest wymaganiem wobec procesu montażu, a nie parametrem odporności wyrobu. Jeżeli złącze jest dostępne dla użytkownika, ochronę na poziomie systemowym trzeba dodać na płycie – niezależnie od tego, jaką wartość HBM podano dla układu.

Nota Texas Instruments SSZB130 wskazuje przy tym, że napięcie ograniczania elementu ochronnego przy wyładowaniu 8 kV według IEC najlepiej przybliżyć pomiar metodą impulsu linii transmisyjnej (TLP) przy prądzie 16 A – czyli przy wartości odpowiadającej przebiegowi normatywnemu w chwili 30 ns po szczycie. Do tego wątku wracamy w drugiej części, przy doborze diod ESD. Metoda TLP polega na podaniu na badany element krótkiego impulsu prądowego o prostokątnym kształcie, wytworzonego przez rozładowanie odcinka linii transmisyjnej, i zmierzeniu napięcia, jakie się na tym elemencie ustala. Powtarzając pomiar przy rosnących wartościach prądu, odtwarza się punkt po punkcie charakterystykę elementu w zakresie prądów udarowych – czyli tam, gdzie zwykła charakterystyka statyczna z karty katalogowej już nie obowiązuje. Do metody wracamy w trzeciej części.

jest przypięta do masy badanego układu, bo przewód masowy sondy tworzy wtedy własną pętlę i sam staje się częścią mierzonego obwodu. Metodyka pomiaru takich zaburzeń to osobny temat, do którego wracamy w trzeciej części.

Systematyczne omówienie wszystkich trzech mechanizmów, wraz z rachunkami dla typowych geometrii, zawiera monografia Henry'ego W. Otta wydana w 2009 roku – rozszerzona wersja jego wcześniejszej książki o redukcji szumów, uzupełniona o rozdziały poświęcone wyładowaniom elektrostatycznym. Dla samego ESD odpowiednikiem jest trzecie wydanie książki Michela Mardiguiana, w którym sprzężenie wyładowania do płytek i wiązek kablowych zostało policzone, z podaniem amplitudy i widma pola w funkcji odległości od miejsca wyładowania.

Punkt sporny: jak daleko sięga wyładowanie elektrostatyczne

Praktyka warsztatowa przyjmuje, że ochrona przed ESD sprowadza się do maty, opaski i uziemionego stanowiska – czyli że zagrożenie istnieje tam, gdzie następuje kontakt. Douglas C. Smith kwestionuje to założenie na podstawie pomiarów. Z jego badań wynika, że prądy wyindukowane w kablu o długości jednego metra przez wyładowanie oddalone o kilka metrów osiągają wartości wystarczające do uszkodzenia współczesnych komponentów, co oznacza, że kontrola ESD ograniczona do samego stanowiska montażowego jest niewystarczająca.

Konsekwencja praktyczna jest niewygodna: wyładowanie do metalowej szafy po drugiej stronie pomieszczenia, do wózka albo do krzesła może wywołać zakłócenie pracy urządzenia, którego nikt nie dotknął. Diagnoza takich zdarzeń jest trudna, bo zdarzenie nie ma świadka i nie zostawia śladu, a objaw – pojedyncze zawieszenie mikrokontrolera, przekłamanie w transmisji – bywa przypisywany oprogramowaniu.

Teza nie jest jednak odosobniona ani lekceważona w normalizacji. Sama norma IEC 61000-4-2 przewiduje, obok wyładowania bezpośredniego do badanego wyrobu, wyładowanie pośrednie do poziomej i pionowej płaszczyzny sprzęgającej, umieszczonych w zdefiniowanej odległości – czyli uznaje istnienie zagrożenia bezkontaktowego. Spór dotyczy raczej zasięgu i tego, czy odległości przyjęte w procedurze badawczej odwzorowują warunki rzeczywistej instalacji.

Smith przepracował dwadzieścia sześć lat w AT&T Bell Laboratories, odchodząc ze stanowiska Distinguished Member of Technical Staff, a następnie kierował rozwojem i badaniami EMC w Auspex Systems; prowadzi zajęcia na Uniwersytecie Oksfordzkim i publikuje materiały techniczne w serwisie emcesd.com. Jego książka o pomiarach wielkiej częstotliwości i szumach w układach elektronicznych z 1993 roku pozostaje podstawową pozycją o metodyce pomiaru zaburzeń impulsowych – wracamy do niej w trzeciej części, przy weryfikacji ochrony.

Przebiegi normatywne to modele, a nie zdarzenia

Najczęstsze nieporozumienie dotyczące badań odporności brzmi: „urządzenie zostało zbadane napięciem 2 kV, więc wytrzyma 2 kV na zaciskach”. Zdanie to jest nieprawdziwe i warto rozebrać dlaczego, bo z tego rozbiór wynika sposób czytania protokołów badań.

Generator używany w badaniu według IEC 61000-4-5 nazywa się generatorem fali kombinowanej. Wartości jego elementów – rezystorów, indukcyjności i kondensatora – dobrano tak, aby przy rozwartych zaciskach dawał falę napięciową 1,2/50 μ s, a przy zwartych – falę prądową 8/20 μ s. Stosunek napięcia przy rozwarciu do prądu przy zwarcu definiuje efektywną impedancję wyjściową generatora, wynoszącą dla portów zasilania 2 Ω . Warto przy tym pamiętać, że w normach kompatybilności elektromagnetycznej słowo port nie oznacza złącza, lecz miejsce styku urządzenia ze środowiskiem zewnętrznym; osobno rozpatruje się port zasilania, port sygnałowy, port uzimienia i port obudowy, a każdy z nich bada się inaczej.

Wynika stąd rzecz zasadnicza: to jest jeden i ten sam impuls. Dla badanego obiektu o dużej impedancji wejściowej jest on próbą napięciową, a dla obiektu o małej impedancji – próbą prądową. Norma stwierdza to wprost i dodaje uwagę o dalszych skutkach: przebieg napięcia i prądu na zaciskach zależy od impedancji wejściowej badanego obiektu, a impedancja ta zmienia się w trakcie samego udaru – na skutek zadziałania zainstalowanych elementów ochronnych albo, gdy ich nie ma lub są niesprawne, na skutek przeskoku iskry lub przebicia elementu. Oznacza to, że „2 kV” z protokołu opisuje nastawę generatora, a nie napięcie, które wystąpiło na wejściu układu.

Drużga rzecz to impedancja źródła, która nie jest jedna. Sieć sprzęgająco-odsprzęgająca to układ pośredniczący między generatorem a badanym urządzeniem: wprowadza udar do wskazanej żyły i jednocześnie odcina go od sieci zasilającej, żeby nie trafił do innych urządzeń w laboratorium. Ma przy tym własną impedancję, którą dodaje do impedancji generatora. Dla krótkich portów sygnałowych sieć ta podnosi efektywną impedancję źródła do 12 Ω przy porcie niesymetrycznym i do 42 Ω przy porcie symetrycznym. Ta sama nastawa napięcia oznacza więc dla portu sygnałowego prąd kilkakrotnie mniejszy niż dla portu zasilania – co bezpośrednio przekłada się na wymaganą wytrzymałość udarową elementu ochronnego. Norma ostrzega przy tym, że dołączenie rezystora podnoszącego impedancję z 2 Ω do na przykład 42 Ω może istotnie zmienić czas narastania i czas do połowy wartości impulsu mierzonego na wyjściu sieci sprzęgającej. Innymi słowy: fala 1,2/50 μ s na porcie sygnałowym nie jest tą samą falą co na porcie zasilania, mimo identycznego oznaczenia.

Dla portów telekomunikacyjnych, czyli wyprowadzonych poza obiekt, stosuje się przebieg zupełnie inny – falę napięciową 10/700 μ s, prądowo 5/320 μ s, o impedancji źródła 40 Ω i poziomach do 4 kV. Jej rodowód wywodzi się z zaleceń ITU-T serii K, w szczególności K.20 i K.21, i odwzorowuje udar piorunowy stłumiony przez kilometry linii, a nie zdarzenie wewnątrz budynku.

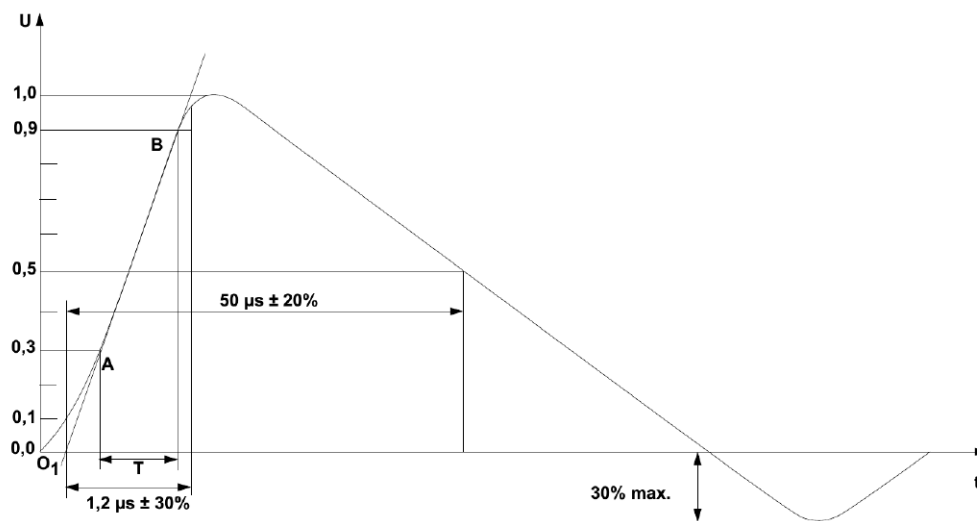
Warto też znać genezę samych kształtów, bo tłumaczy ona, dlaczego są takie, a nie inne. Definicje przebiegów w IEC 61000-4-5 wywodzą się z normy IEC 60060-1, dotyczącej technik badań wysokonapięciowych. W klasycznej technice wysokich napięć impulsy napięciowy i prądowy bada się jednak osobno: generator 1,2/50 μ s służył do prób izolacji i oddawał do obciążenia o dużej impedancji prądy rzędu miliamperów, a generator 8/20 μ s służył do prób ograniczników przepięć i oddawał duży prąd do obciążenia o małej impedancji. Połączenie obu funkcji w jednym urządzeniu jest wynalazkiem norm kompatybilności elektromagnetycznej, podyktowanym tym, że badany wyrób nie jest ani czystą izolacją, ani czystym ogranicznikiem.

Trzy rachunki, które warto wykonać przed doбором elementu

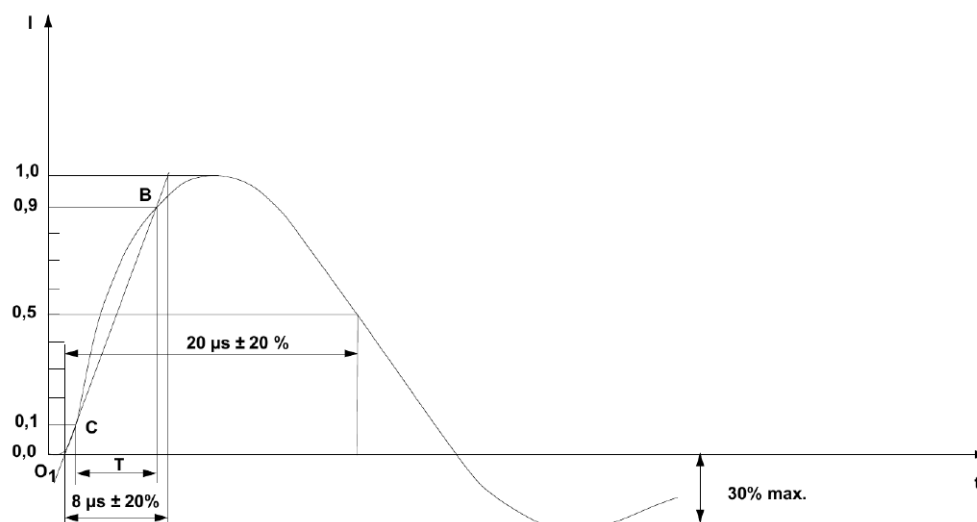
Dobór elementu ochronnego zaczyna się od trzech wielkości: prądu, jaki popłynie, energii, jaka odłoży się w elemencie, oraz napięcia, jakie mimo wszystko dotrze do chronionego układu. Każdą z nich da się oszacować w jednym wierszu, a wyniki bywają odległe od intuicji.

Rachunek pierwszy: jaki prąd popłynie

Prąd udarowy wynika z napięcia próby podzielonego przez sumę impedancji generatora i obwodu.



a) napięcie przy rozwartych zaciskach generatora



b) prąd przy zwartych zaciskach generatora

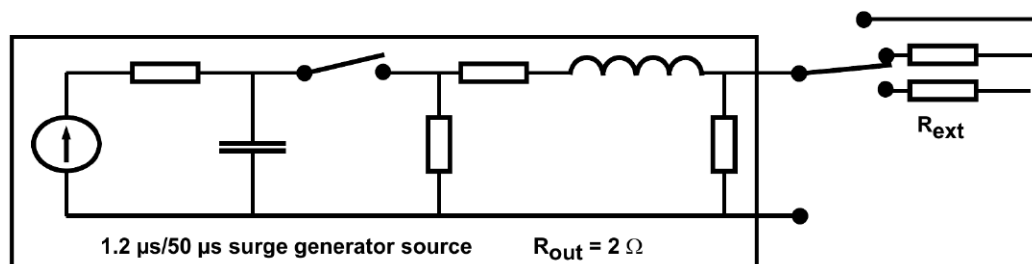
Rysunek 2a, 2b. Przebiegi definiujące falę kombinowaną, z naniesionymi punktami odniesienia 0,1, 0,3, 0,5 i 0,9 wartości szczytowej oraz tolerancjami czasu czoła i czasu do połowy wartości. Podstawa: STMicroelectronics, nota aplikacyjna AN4275 „IEC 61000-4-5 standard overview”, DocID024389 Rev 1, sierpień 2013, rysunki 2 i 3

$$I_{sz} = \frac{U_n}{Z_{gen} + Z_{obw}}$$

Przy udarze 1 kV w obwodzie o impedancji 2 Ω daje to 500 A. Liczba ta pochodzi z materiałów projektowych firmy Bourns dotyczących ochrony portu RS-485 i prowadzi tam do konkretnego wniosku: tyrystorowy ochronnik o wytrzymałości 300 A w przebiegu 8/20 μs jest w tych warunkach za słaby i trzeba sięgnąć po wersję o wytrzymałości 800 A.

Rachunku nie trzeba jednak wykonywać dla każdego przypadku osobno, bo wyniki dla wszystkich kombinacji poziomu badania i rezystora zewnętrznego są stabelaryzowane. Poniższe zestawienie pochodzi z noty aplikacyjnej STMicroelectronics AN4275 i podaje maksymalny prąd szczytowy przy zwartym wyjściu generatora.

Z tabeli 2 wynika reguła doborowa istotniejsza niż sam wynik pojedynczego rachunku: przy tej samej nastawie napięcia port zasilania i symetryczny port sygnałowy



Rysunek 3. Układ generatora z zaznaczeniem elementów decydujących o obu kształtach przebiegu oraz rezystora zewnętrznego R_{ext} ustalającego impedancję źródła. Podstawa: STMicroelectronics AN4275, rysunek 1

Tabela 2. Maksymalny prąd szczytowy przy zwartym wyjściu generatora, w zależności od poziomu badania i rezystora zewnętrznego ustalającego impedancję źródła. Źródło: STMicroelectronics, nota aplikacyjna AN4275, tabela 2

Poziom badania	$R_{eq} = 2 \Omega$ (zasilanie)	$R_{eq} = 12 \Omega$ (port niesym.)	$R_{eq} = 42 \Omega$ (port sym.)
25 V – klasa 0	12,5 A	2,1 A	0,6 A
500 V – klasa 1	250 A	42 A	12 A
1 kV – klasa 2	500 A	84 A	24 A
2 kV – klasa 3	1000 A	167 A	48 A
4 kV – klasa 4	2000 A	334 A	96 A

różnią się prądem udarowym ponad dwudziestokrotnie. Elementu dobraneo dla portu sygnałowego nie wolno więc przenosić na port zasilania bez powtórzenia obliczenia – i odwrotnie, dobieranie do portu sygnałowego elementu o wytrzymałości wymaganej na zasilaniu oznacza niepotrzebnie dużą pojemność wnoszoną do toru.

Rachunek drugi: ile energii odłoży się w elemencie

Energia jest iloczynem napięcia ograniczania i ładunku, który przez element przepłynie; ładunek dla przebiegu 8/20 μ s można przybliżyć trójkątem o wysokości równej prądowi szczytowemu i podstawie rzędu 30 μ s.

$$E \approx U_{ogr} \cdot Q \quad Q \approx \frac{1}{2} I_{sz} \tau_{ef}$$

Dla policzonych wyżej 500 A ładunek wynosi około 7,5 mC, a przy napięciu ograniczania rzędu 40 V energia odkładana w elemencie osiąga około 0,3 J. Dla porównania: wyładowanie elektrostatyczne 8 kV według IEC 61000-4-2 niesie około 4,8 mJ, czyli mniej więcej sześćdziesiąt razy mniej. To jest liczbową treść stwierdzenia, że dioda dobrana do ochrony przed ESD nie jest elementem do ochrony przed udarem – nie dlatego, że ma złe napięcie ograniczania, tylko dlatego, że nie ma dokąd odprowadzić dwóch rzędów wielkości więcej energii.

Rachunek trzeci: ile napięcia i tak przejdzie

Ostatni rachunek jest najbardziej niewygodny, bo dotyczy wielkości, której nie ma w żadnej karcie katalogowej elementu ochronnego – indukcyjności doprowadzeń.

$$U_L = L \cdot \frac{di}{dt}$$

Odcinek ścieżki o długości kilkunastu milimetrów wraz z przelotką ma indukcyjność rzędu 10 nH. Przy wyładowaniu elektrostatycznym o prądzie szczytowym 30 A narastającym w ciągu 1 ns szybkość zmian prądu wynosi $3 \cdot 10^{10}$ A/s, a spadek napięcia na tej indukcyjności – około 300 V. Napięcie ograniczania samej diody wynosi w takim zastosowaniu kilkanaście woltów.

Wniosek jest jednoznaczny i wart zapamiętania na całą resztę cyklu: przy zaburzeniach o krótkim czasie narastania o napięciu na chronionym wejściu decyduje nie parametr elementu ochronnego, lecz geometria drogi, jaką prąd zaburzenia płynie do punktu odniesienia. Element o napięciu ograniczania niższym o dwa woltu nie poprawi niczego, jeżeli droga powrotna jest dłuższa o centymetr. Do konsekwencji projektowych wracamy w trzeciej części cyklu.

Transient Control Level – dlaczego pracujemy w trybie napraw

Pozostaje pytanie, dlaczego konstruktor urządzenia w ogóle musi wykonywać powyższe rachunki samodzielnie. W energetyce wysokiego napięcia analogiczny problem rozwiązano dawno, wprowadzając poziomy izolacji znamionowej: dla danego napięcia sieci ustala się poziom, którego przepięcia nie przekroczą dzięki zastosowanym ogranicznikom, a urządzenia projektuje się na wytrzymałość wyższą od tego poziomu. Skutek jest taki, że sprzęt zaprojektowany zgodnie z normami nie ulega uszkodzeniom od udarów piorunowych.

Sylwetki

François D. Martzloff

Inżynier elektroniki, którego kariera objęła Southern States Equipment, dział badawczo-rozwojowy General Electric oraz National Institute of Standards and Technology. Uznawany za jedną z osób, które ukształtowały dorobek w dziedzinie urządzeń i norm ochrony przed przepięciami. Autor ośmioczęściowej antologii wydanej przez NIST, poświęconej ochronie przepięciowej w obwodach niskiego napięcia prądu przemiennego, obejmującej bibliografię komentowaną oraz przedruki prac własnych i innych autorów. W 2012 roku otrzymał od IEEE Standards Association nagrodę za całokształt dorobku – za rzetelność, przywództwo i mentorstwo w rozwoju norm dotyczących urządzeń ochrony przepięciowej. Prace wykorzystane w tej części: zależność częstości wystąpień od poziomu napięcia, koordynacja ochronników w obwodach niskiego napięcia oraz rozptyw prądu piorunowego po bezpośrednim wyładowaniu w budynek.

Ronald B. Standler

Autor monografii Protection of Electronic Circuits from Overvoltages, wydanej przez Wiley w 1989 roku i wznowionej przez Dover. Książka powstała jako odpowiedź na rozproszenie literatury przedmiotu, która przed jej publikacją była porozrzucana po czasopismach, patentach, materiałach konferencyjnych i raportach wojskowych. Obejmuje uszkodzenia i zakłócenia pracy, zagrożenia środowiskowe, standardowe przebiegi laboratoryjne, właściwości nieliniowych elementów ochronnych, zastosowania w torach sygnałowych, zasilaczach prądu stałego i sieci niskiego napięcia do 1 kV, a także weryfikację przyjętych rozwiązań – wraz z rozdziałem poświęconym indukcyjności pasożytniczej, do którego wracamy w drugiej i trzeciej części cyklu.

W 1976 roku F.A. Fisher i F.D. Martzloff zaproponowali przeniesienie tej samej zasady do instalacji niskiego napięcia pod nazwą Transient Control Level. Propozycja obejmowała ustalenie zestawu znormalizowanych poziomów kontroli przepięć i takie skoordynowanie ograniczników z wytrzymałością urządzeń, aby oba parametry były wobec siebie zdefiniowane, a nie dobierane niezależnie przez producenta instalacji i producenta wyrobu.

Propozycja nie została wdrożona. Autorzy wskazali powód: sprzęt trafiał już na rynek bez takiej koordynacji, więc stan faktyczny wyprzedził normalizację. Konsekwencję, którą opisali, mamy do dziś – uszkodzenia zdarzają się,

a środki zaradcze wprowadza się po fakcie, jako modyfikacje konstrukcji już wyprodukowanej.

Dla czytelnika projektującego urządzenie oznacza to rzecz konkretną. Nie istnieje wartość napięcia, którą można przyjąć jako „poziom, do którego wystarczy projektować”, ponieważ nikt takiego poziomu nie gwarantuje. Poziom trzeba ustalić samodzielnie – na podstawie klasy ekspozycji instalacji, w której urządzenie będzie pracować, oraz założeń przyjętych przez projektanta ochrony instalacyjnej, jeżeli takowa istnieje. To jest właśnie ta część pracy, której nie ma w karcie katalogowej, i to jest powód, dla którego cały pierwszy artykuł cyklu dotyczy zagrożenia, a nie elementów.

Słowniczek

Zbiórce zestawienie terminów objaśnianych w tej części cyklu, w ramach i w tekście. Terminy wprowadzane w częściach drugiej i trzeciej zostaną dopisane do słowniczka w tamtych artykułach.

CDM, model naładowanego elementu – Model wyładowania, w którym ładunek gromadzi się w obudowie układu scalonego i rozładowuje przez wyprowadzenie stykające się z uziemioną powierzchnią. Czas narastania poniżej nanosekundy. Opisany w ANSI/ESDA/JEDEC JS-002.

Czas czoła – Umowny czas narastania impulsu, pierwsza z dwóch liczb w zapisie typu 8/20 μ s. Wyznaczany przez ekstrapolację prostej poprowadzonej przez dwa punkty odniesienia na zboczu narastającym.

Czas do połowy wartości – Druga liczba w zapisie typu 8/20 μ s: czas od umownego początku impulsu do chwili, w której wielkość opada do połowy wartości szczytowej.

CDN, sieć sprzęgająco-odsprężająca – Układ między generatorem udaru a badanym urządzeniem. Wprowadza udar do wskazanej żyły i odcina go od sieci zasilającej. Wnosi własną impedancję, która sumuje się z impedancją generatora.

Dioda TVS – Dioda ograniczająca przepięcia (transient voltage suppressor), pracująca w kierunku zaporowym i przewodząca po przekroczeniu napięcia przebicia lawinowego. Omawiana szczegółowo w drugiej części cyklu.

Energia właściwa – Całka kwadratu prądu po czasie, wyrażana w MJ/ Ω . Miara ciepła, jakie prąd wydzieliłby w rezystancji jednego oma, a więc miara obciążenia cieplnego elementu.

ESD – Wyładowanie elektrostatyczne. Rozróżnia się modele komponentowe (HBM, CDM), opisujące montaż i transport układów scalonych, oraz model systemowy z IEC 61000-4-2, opisujący gotowy wyrób w rękach użytkownika.

Fala kombinowana – Przebieg wytwarzany przez generator, który przy rozwartych zaciskach daje falę napięciową 1,2/50 μ s, a przy zwartych falę prądową 8/20 μ s. Stosunek obu wartości szczytowych definiuje efektywną impedancję wyjściową generatora.

GDT, odgromnik gazowy – Element przelączający napięcie: po przekroczeniu napięcia zapłonu przechodzi w łuk o bardzo małej rezystancji i zwiera obwód. Omawiany w drugiej części cyklu.

HBM, model ciała ludzkiego – Model wyładowania odwzorowujący rozładowanie naładowanego człowieka: kondensator 100 pF przez rezystor 1500 Ω . Opisany w ANSI/ESDA/JEDEC JS-001. Dotyczy warunków montażu, nie odporności gotowego wyrobu.

Impedancja źródła – Iloraz napięcia przy rozwarciu i prądu przy zwarcu na wyjściu generatora. Decyduje o tym, jaki prąd popłynie przy zadanej nastawie napięcia: 2 Ω dla portów zasilania, 12 Ω i 42 Ω dla portów sygnałowych, 40 Ω dla telekomunikacyjnych.

Klasa ekspozycji – Formalny opis środowiska instalacyjnego, z którym norma wiąże poziom badania. Klasy 0...5 według IEC 61000-4-5, od pomieszczenia wydzielonego po sprzęt przyłączony do linii napowietrznych.

LPL, poziom ochrony odgromowej – Jeden z czterech

poziomów zdefiniowanych w IEC 62305, określających zakładane parametry prądu piorunowego. LPL I jest najostrzejszy: prąd szczytowy 200 kA, ładunek 100 C.

Napięcie ograniczania – Napięcie, jakie ustala się na elemencie ochronnym przy przepływie zadanego prądu udarowego. Nie jest stałą elementu – zależy od prądu i od rezystancji dynamicznej.

PEN – Wspólny przewód pełniący jednocześnie funkcję neutralną i ochronną w układzie TN-C. Jego przerwanie prowadzi do przepięcia dorywczego u odbiorcy.

Port – W normach kompatybilności elektromagnetycznej: miejsce styku urządzenia ze środowiskiem zewnętrznym, a nie złącze. Osobno rozpatruje się port zasilania, sygnałowy, uziemienia i obudowy.

Proces przejściowy – Zdarzenie jednorazowe o czasie trwania od pojedynczych nanosekund do setek mikrosekund: wyładowanie elektrostatyczne, paczka szybkich procesów przejściowych, udar łączeniowy lub piorunowy.

Przepięcie dorywcze (TOV) – Podwyższone napięcie trwające od kilku okresów sieci do godzin. Nie jest procesem przejściowym i to ono, a nie udar mikrosekundowy, najczęściej niszczy warystory.

Przepięcie różnicowe i wspólne – Różnicowe występuje między przewodami toru i jest widoczne na wejściu układu wprost. Wspólne występuje między wszystkimi przewodami łącznie a punktem odniesienia i nie jest widoczne pomiędzy przewodami sygnałowymi.

Przetężenie – Nadmierny prąd, nie napięcie. Wymaga innych środków niż przepięcie: bezpiecznika, elementu PTC albo ogranicznika elektronicznego.

Prąd następczy – Prąd płynący przez element przelączający napięcie po ustaniu udaru, podtrzymywany już samym napięciem roboczym obwodu. Problem charakterystyczny dla odgromników gazowych, omawiany w drugiej części.

ŚN, średnie napięcie – Napięcia powyżej 1 kV do 30 lub 35 kV według norm polskich. Poziom linii zasilających transformatory osiedlowe.

TCL, Transient Control Level – Koncepcja Fishera i Martzloffa z 1976 roku: hierarchia znormalizowanych poziomów kontroli przepięć, wiążąca parametry ograniczników z wytrzymałością urządzeń. Nigdy nie wdrożona w instalacjach niskiego napięcia.

TLP, pomiar impulsem linii transmisyjnej – Metoda badania elementów ochronnych prostokątnym impulsem prądowym z odcinka linii transmisyjnej. Pozwala odtworzyć charakterystykę elementu w zakresie prądów udarowych.

TN-C – Układ sieci, w którym funkcje neutralna i ochronna są połączone w jednym przewodzie PEN.

Warystor (MOV) – Element ograniczający napięcie o silnie nieliniowej charakterystyce, wykonany ze spieku tlenku cynku. Degraduje się z każdym udarem. Omawiany w drugiej części cyklu.

Tabela 3. Zestawienie dokumentów normalizacyjnych przywoływanych w trzech częściach cyklu

Dokument	Czego dotyczy	Gdzie wraca w cyklu
IEC 61000-4-2	Odporność na wyładowania elektrostatyczne; wyładowanie kontaktowe i powietrzne, bezpośrednie i pośrednie do płaszczyzn sprzęgających	część 1 i 2 – dobór diod ESD
IEC 61000-4-4	Odporność na szybkie procesy przejściowe seryjne; sprzężenie pojemnościowe do wiązek kablowych	część 2 – dławiki i filtry
IEC 61000-4-5	Odporność na udar; generator fali kombinowanej, sieci sprzęgająco-odsprzęgające	część 2 i 3 – dobór i koordynacja
IEC 60060-1	Techniki badań wysokonapięciowych; rodowód kształtów 1,2/50 μ s i 8/20 μ s	część 1 – geneza przebiegów
IEC 62305	Ochrona odgromowa; poziomy ochrony LPL, parametry prądu piorunowego	część 1 i 3 – warunki brzegowe
EN 61643-11	Ograniczniki przepięć do sieci niskiego napięcia; typy 1, 2 i 3	część 3 – wejście sieciowe
ITU-T K.20, K.21, K.44	Odporność portów telekomunikacyjnych; fala 10/700 μ s, koordynacja z rezystorem szeregowym	część 3 – linie zewnętrzne
ANSI/ESDA/JEDEC JS-001, JS-002	Modele komponentowe HBM i CDM; procedury badań układów scalonych	część 1 – rozgraniczenie poziomów
ANSI/ESD STM5.5	Pomiar metodą impulsu linii transmisyjnej (TLP)	część 3 – weryfikacja
UL 1449	Ograniczniki przepięć; badanie przy przepięciu dorywczym o ograniczonym prądzie	część 2 – starzenie warystorów

Trzy pytania przed pierwszym elementem na schemacie

Martzloff formułował zadanie projektanta jako zestaw pytań, na które trzeba odpowiedzieć, zanim dobierze się jakiegokolwiek urządzenie ochronne na styku instalacji i sprzętu. W wersji dostosowanej do projektowania urządzenia, a nie obiektu, brzmią one następująco.

1. Jakie jest zagrożenie na tym konkretnym porcie? Odpowiedź musi zawierać cztery liczby: kształt przebiegu, amplitudę, impedancję źródła i szacowaną częstość wystąpień. Port zasilania, port sygnałowy wewnątrzobiektywny i port wyprowadzony poza budynek dają trzy różne odpowiedzi i trzy różne rozwiązania.
2. Jaka jest wytrzymałość chronionego układu? Nie ta z karty katalogowej dotycząca modelu HBM, tylko

rzeczywista wartość napięcia i prądu, przy której wejście ulega uszkodzeniu – z uwzględnieniem tego, że przez wewnętrzne diody ochronne układu prąd nie może płynąć w nieskończoność.

3. Jaki poziom pośredni chcemy wymusić i czym? To jest pytanie o architekturę ochrony: jeden element czy kaskada, gdzie umieszczona, z jaką impedancją rozdzielającą między stopniami. Odpowiedzi na nie szukamy w częściach drugiej i trzeciej.

Druga część cyklu zajmie się elementami: fizyką ich działania, granicami stosowności, mechanizmami starzenia i kompromisami, których nie da się ominąć – od pojemności diod niskopojemnościowych, przez ucieczkę termiczną warystorów, po prąd następczy odgromników gazowych.

JT, WM

Bibliografia

1. R.B. Standler, Protection of Electronic Circuits from Overvoltages, Wiley 1989/Dover 2002.
2. F.D. Martzloff, Surge Protection in Low-Voltage AC Power Circuits – An 8-part Anthology, NISTIR 6714, NIST: <https://www.nist.gov/pml/quantum-measurement/surge-protection-low-voltage-ac-power-circuits/spd-anthology-annotated/spd>
3. F.D. Martzloff, Surge Protection Techniques in Low-Voltage AC Power Systems, NIST – podstawa rys. 1.
4. F.A. Fisher, F.D. Martzloff, praca o Transient Control Levels, antologia NIST cz. 2: <https://www.nist.gov/document/allpaperspart2pdf>
5. F.D. Martzloff, On the Dispersion of Lightning Current after a Direct Flash to a Building, ICLP, Rodos 2000.
6. F.D. Martzloff, Surge Protection of End-User Equipment, NIST.
7. W. Meissen, Überspannungen in Niederspannungsnetzen, ETZ 104, 1983.
8. P. Hasse, Overvoltage Protection of Low Voltage Systems, wyd. 2, IET 2000.
9. H.W. Ott, Electromagnetic Compatibility Engineering, Wiley 2009.
10. M. Mardiguian, Electro Static Discharge: Understand, Simulate, and Fix ESD Problems, wyd. 3, Wiley-IEEE 2009.
11. D.C. Smith, High Frequency Measurements and Noise in Electronic Circuits, Van Nostrand Reinhold 1993; materiały techniczne w serwisie emcesd.com.
12. C. Duvvury, H. Gossner (red.), System Level ESD Co-Design, Wiley 2015.
13. S.H. Voldman, ESD: Physics and Devices oraz ESD Basics, Wiley.
14. Texas Instruments, System-Level ESD Protection Guide, SSZB130.
15. STMicroelectronics, IEC 61000-4-5 standard overview, nota aplikacyjna AN4275, DocID024389 Rev 1, 2013 – podstawa rys. 2a, 2b, 3 i tab. 2.
16. Bourns, materiały projektowe do ochrony portu RS-485 (TBU-RS, TISP).
17. H. Zenkner, TVS Diodes: Basics, Function, and Applications, nota aplikacyjna ANP143, Würth Elektronik: <http://www.we-online.com/anp143>
18. Normy i zalecenia: IEC 61000-4-2, IEC 61000-4-4, IEC 61000-4-5, IEC 60060-1, IEC 62305, EN 61643-11, UL 1449, ITU-T K.20, K.21 i K.44, ANSI/ESDA/JEDEC JS-001 i JS-002, ANSI/ESD STM5.5.

Czujniki wilgotności i cieczy dla domu i ogrodu

Część A – projekt przykładowy: czujnik wilgotności gleby z sondą grzebieniową i wyjściem przekąźnikowym

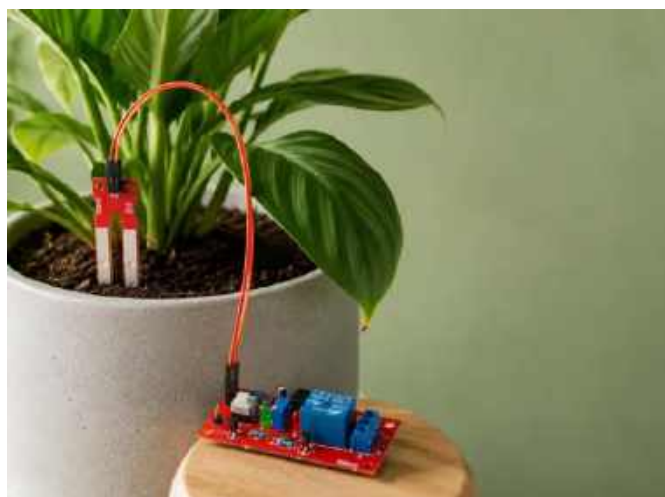
Wilgotność bywa cichym sprzymierzeńcem i cichym sabotażystą naraz: za mało wody w doniczce – roślina usycha, za dużo na podłodze piwnicy – grozi zalaniem i porażeniem. W praktyce domowej i ogrodowej mierzy się trzy różne rzeczy: wilgotność powietrza, wilgotność gleby oraz samą obecność cieczy. Najprostsze, a zarazem bardzo użyteczne, są układy progowe – takie, które nie podają liczbowej wartości, lecz sygnalizują przekroczenie granicy i mogą od razu coś załączyć: pompę, zawór, alarm. Omówienie tego tematu zawiera dwie części. W części A prezentujemy projekt przykładowy – progowy czujnik wilgotności gleby z sondą grzebieniową i wyjściem przekąźnikowym. W części B poddajemy go komentarzowi redakcyjnemu, omawiamy spektrum metod pomiaru wilgotności i wykrywania cieczy oraz zestawiamy projekty pokrewne i ofertę kitów AVT.

Część A. Projekt przykładowy Czujnik wilgotności gleby z sondą grzebieniową (LM393)

Pomiar wilgotności gleby przydaje się nie tylko przy uprawach. Bywa też sposobem na wykrycie wody tam, gdzie jej być nie powinno – na przykład wokół instalacji elektrycznej, gdzie wilgotne podłoże stwarza zagrożenie. Obok klasycznego higrometru sprawdzają się właśnie proste układy progowe: dają sygnalizację wizualną i w razie potrzeby uruchamiają lokalne działanie, gdy wilgotność wykryta przez sondę wykroczy poza dopuszczalny zakres. Prezentowany układ to swego rodzaju higrostat: właściwym elementem czułym jest grzebieniowa sonda z PCB, a dołączony do niej układ z komparatorem i przekąźnikiem rozróżnia wykryty stan.

Zasada działania

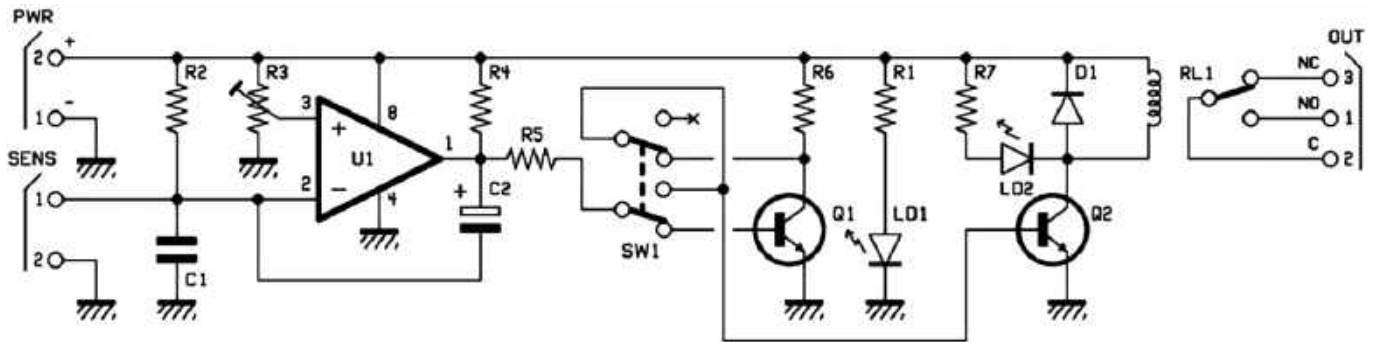
Sonda grzebieniowa to w istocie dwie odizolowane od siebie elektrody, dołączone dwużyłowym przewodem do zacisków SENS. Jedna trafia do masy, druga – razem z rezystorem R2 dołączonym do linii zasilania – tworzy dzielnik napięcia. Po wsunięciu sondy w badane podłoże między jej stykami pojawia się rezystancja: im wilgotniejsza gleba, tym rezystancja mniejsza,



Projekt: progowy czujnik wilgotności gleby (i detektor obecności wody) z sondą grzebieniową wykonaną z płytki drukowanej, komparatorem LM393 i wyjściem przekąźnikowym z sygnalizacją LED. Podwójny przełącznik SW1 pozwala wybrać, czy przekąźnik ma zadziałać po przekroczeniu, czy po spadku progu wilgotności.

Źródło: Boris Landoni, „Soil Moisture Sensor”, Open-Electronics.org. Pliki Gerber do pobrania z materiału źródłowego.

Zastosowanie: monitoring wilgotności gleby, automatyczne podlewanie (pompa lub zawór elektromagnetyczny), a także wykrywanie zalania lub obecności wody na podłodze (alarm, pompa odsysająca).



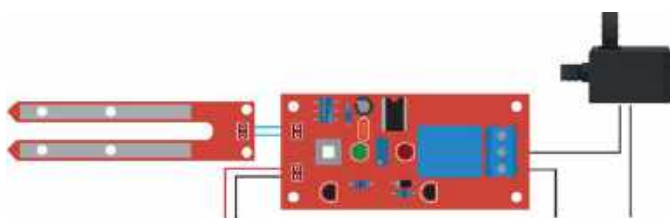
Rysunek 1. Schemat ideowy czujnika: komparator LM393 z pojemnościowym sprzężeniem zwrotnym, stopień przekaźnikowy i sonda grzebieniowa

a napięcie na wejściu odwracającym komparatora niższe. Równolegle do zacisków sondy dołączono kondensator filtrujący zakłócenia impulsowe, które mogłyby zaburzyć pracę komparatora.

Układ zbudowano na kostce LM393 zawierającej dwa komparatory z wyjściami typu otwarty kolektor (stąd rezystor R4 podciągający wyjście na wyprowadzeniu 1); wykorzystano jeden z nich. Komparator U1 połączono nietypowo – ze sprzężeniem zwrotnym przez kondensator (C2), przez co pracuje jak jednozasilaniowy integrator progowy: jego napięcie wyjściowe narasta zależnie od napięcia wejściowego, a poziom odniesienia ustala się trimmerem R3 na wejściu nieodwracającym. Dzięki temu dla danej wilgotności wyjście dąży płynnie do stanu ustalonego, a nie przełącza się skokowo.

Wyjście przekaźnikowe i wybór progu

Za komparem rezystor R5 doprowadza jego stan do podwójnego przełącznika SW1, który decyduje o kierunku działania. W pierwszym położeniu (spoczynkowym, jak na schemacie) wyjście U1 steruje tranzystorem Q1, a ten – poprzez Q2 – cewką przekaźnika RL1. W tym trybie przekaźnik zadziała, gdy wilgotność jest zbyt mała (gleba przesuszona lub brak wody): styk C–NO może wtedy załączyć pompę podlewającą albo zawór elektromagnetyczny. W drugim położeniu SW1 wyjście komparatora steruje bezpośrednio tranzystorem Q2 (Q1 jest wyłączony z obwodu), a przekaźnik zadziała przy nadmiarze wilgoci – na przykład gdy na podłodze pojawi się woda: styk może uruchomić pompę odsysającą lub alarm, albo poprzez C–NC odciąć zasilanie zraszacza.

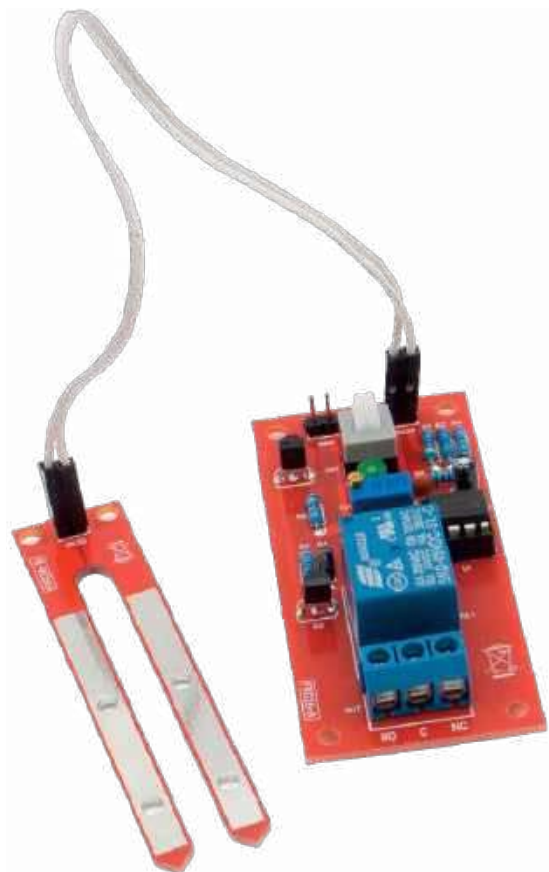


Rysunek 2. Podłączenie układu do sterowania pompą normalnie wyłączoną, która załącza się po wystąpieniu stanu alarmowego

Zasilanie doprowadza się do zacisków PWR: wymagane jest dobrze stabilizowane 5 V – to warunek stabilnych progów przełączania. Dioda LD1 (przez rezystor R1) sygnalizuje obecność zasilania, a LD2 (przez R7), włączona równolegle do cewki, sygnalizuje zadziałanie przekaźnika. Antyrównolegle do cewki dołączono diodę D1, która gasi przepięcie indukcyjne przy wyłączaniu Q2 i chroni go przed uszkodzeniem.

Montaż i uruchomienie

Urządzenie tworzą dwie płytki: główna, dwustronna, z elektroniką, oraz sonda w kształcie dwuzębego grzebienia, również wykonana z laminatu (można zastosować inną sondę dwuelektrodową – ważne, by miała dwie



Fotografia Kompletny czujnik: sonda grzebieniowa i płytka z komparatorem oraz przekaźnikiem

długie metalowe elektrody dołączane do zacisków SENS). Montaż rozpoczyna się od rezystorów i diody, potem podstawka LM393 (4+4), trimmer, tranzystory, diody LED, kondensatory, przełącznik SW1 (do druku), a na końcu miniaturowy przekaźnik i trzypozycyjne złącze wyjściowe o rastrze 5 mm. Elementy biegunowe (diody, tranzystory, kondensator elektrolityczny, orientacja układu w podstawce) montuje się zgodnie z rysunkiem montażowym.

Do zacisków sondy przylutowuje się cienki przewód (0,5 mm² lub mniej), którego drugie końce trafiają do SENS 1...2. Układ zasilany jest dobrze stabilizowanym napięciem 5 V przy prądzie ok. 100 mA. Po włączeniu zaświeca się LD1; następnie ustawia się przełącznik SW1 i trimmer R3 tak, aby dla wybranej funkcji przekaźnik zadziałał przy pożądanym progu. Aby mierzyć wilgotność gleby, sondę wbija się w ziemię; aby wykrywać obecność wody – kładzie się ją płasko na podłożu.

Kluczowe elementy

Tabela 1. Najważniejsze elementy układu i ich rola (wartości – wg plików źródłowych)

Element	Rola w układzie
U1 – LM393	podwójny komparator (open-collector); jeden kanał pracuje jako integrator progowy
R2	rezystor dzielnika napięcia współpracujący z rezystancją sondy
R3 (trimmer)	nastawa progów – napięcie odniesienia komparatora
R4	podciąganie wyjścia typu otwarty kolektor
C1/C2	filtr zakłóceń na zaciskach sondy/sprzężenie zwrotne (integrator)
Q1, Q2	tranzystory NPN sterujące cewką przekaźnika
RL1	przekaźnik wykonawczy (styki C/NO/NC)
D1	dioda gasząca przepięcie cewki przekaźnika
LD1/LD2	sygnalizacja zasilania/zadziałania przekaźnika
SW1	podwójny przełącznik wyboru kierunku progów
SENS/PWR/wyjście	złącza: sonda/zasilanie 5 V/styk przekaźnika

Bezpieczeństwo. Jeśli styk przekaźnika przełącza pompę lub zawór zasilane z sieci 230 V, obwód wykonawczy należy odseparować od niskonapięciowej części 5 V, zadbać o odpowiednią izolację i obudowę, a przy obciążeniach indukcyjnych – o ochronę przeciwprzepięciową styków.

Część B. Komentarz, dyskusja rozwiązań i projekty pokrewne

Część B prowadzimy w trzech wątkach. Najpierw podajemy projekt przykładowy komentarzowi redakcyjnemu, następnie omawiamy spektrum metod pomiaru

wilgotności i wykrywania cieczy, a na koniec zestawiamy projekty pokrewne, literaturę oraz ofertę kitów AVT.

Komentarz redakcyjny do projektu

To klasyczny, tani i skuteczny układ progowy: komparator, przekaźnik i sonda z płytki. Jego atutem jest dwukierunkowy próg (SW1) i gotowe wyjście przekaźnikowe do pompy, zaworu lub alarmu. Poniżej zbieramy uwagi i kierunki rozwoju.

Pomiar rezystancyjny to elektroliza i dryf – stała składowa napięcia na sondzie w wilgotnej glebie wywołuje elektrolizę i korozję elektrod, a czasem – dryf odczytu (naszybki dryf skarży się zresztą komentarz pod artykułem źródłowym). Lekarstwem jest zasilanie sondy prądowe lub impulsowe, elektrody odporne (grafit, stal nierdzewna, złoczone) albo – najlepiej – pomiar pojemnościowy, w którym elektrody nie stykają się galwanicznie z glebą.

To detektor progowy, nie miernik – układ pokazuje stan „mokro/sucho”, a nie liczbową wilgotność. Do sterowania podlewaniami to w zupełności wystarcza; do monitoringu czy zbierania danych trzeba przetwornika ADC, kalibracji i najlepiej sondy pojemnościowej z wyjściem napięciowym.

Kalibracja i stabilność progów – próg zależy nie tylko od trimmera R3, ale i od rodzaju gleby, jej zasolenia oraz temperatury, więc wymaga kalibracji „na miejscu”. Autor słusznie podkreśla, że 5 V musi być dobrze stabilizowane – inaczej progi „pływają”.

Trwałość w terenie – czujnik pracuje w wilgoci, więc płytkę warto zabezpieczyć lakierem i obudowę o odpowiednim stopniu IP, a przy przełączaniu pompy sieciowej – zadbać o separację i tłumienie przepięć. To detale, które decydują o tym, czy układ przetrwa sezon w ogrodzie.

Trzy różne zadania – tytuł artykułu obejmuje wilgotność powietrza (RH), wilgotność gleby i obecność cieczy. To trzy różne pomiary i trzy różne klasy czujników – warto je rozróżniać, bo mylenie ich prowadzi do złego doboru sensora.

Jak mierzy się wilgotność i obecność cieczy – przegląd rozwiązań

Czujniki rezystancyjne (konduktometryczne). Dwie elektrody, których rezystancja maleje wraz z wilgotnością – rozwiązanie z projektu przykładowego. Tanie i proste, ale podatne na elektrolizę, korozję i dryf; poprawia je zasilanie prądowe.

Czujniki pojemnościowe. Wilgotność zmienia przenikalność dielektryczną gleby pełniącą rolę dielektryka między okładkami. Elektrody pokryte maską nie stykają się z glebą, więc nie korodują; popularne moduły dają na wyjściu napięcie proporcjonalne do wilgotności. Najczęściej spotyka się tani moduł „soil moisture v1.2/v2.0”,



Internet Rzeczy w pomiarach środowiskowych Zestaw do kontrolowania wilgotności gleby Enviro Grow

Pielęgnacja kwiatów doniczkowych to wdzięczne zajęcie poprawiające atmosferę domu, ale też świetna okazja do poznania układów IoT zawierających nie tylko czujniki, lecz także elementy wykonawcze. Doskonałym do tego celu produktem, umożliwiającym natychmiastowe uruchomienie działającego urządzenia, okazuje się zestaw do kontrolowania wilgotności gleby uprawianych roślin Enviro Grow (PIM637) firmy Pimoroni. Zestaw przeznaczony jest do projektów automatycznego ogródka.

Wydanie EP06/2024 dostępne na
www.UlubionyKiosk.pl

Tabela 2. Porównanie metod pomiaru wilgotności i wykrywania cieczy

Metoda	Zasada	Zalety	Ograniczenia
Rezystancyjna	rezystancja między elektrodami maleje z wilgocą	prosta, tania (jak w projekcie)	elektroliza, korozja, dryf; lepsza z zasilaniem AC
Pojemnościowa	wilgoć zmienia przenikalność dielektryka gleby	bez kontaktu galwanicznego, odporna na korozję	wymaga układu przetwarzania i kalibracji
Higrometr RH (powietrze)	cyfrowy pomiar wilgotności względnej	skalibrowany, cyfrowy (I ² C)	mierzy powietrze, nie glebę
Czujnik zalania	progowe elektrody/taśma/pływak	szybki alarm wycieku	tylko stan tak/nie

w którym generator na układzie 555 (ok. 1,5 MHz) z prostownikiem szczytowym zamienia pojemność sondy na napięcie wyjściowe.

Higrometry wilgotności powietrza (RH). Cyfrowe czujniki, jak DHT22, Sensirion SHT4x czy Bosch BME280, podają skalibrowaną wilgotność względną

(a często i temperaturę) przez interfejs I²C – do pomiaru powietrza, nie gleby.

Czujniki zalania i obecności wody. Proste elektrody progowe, taśmy detekcyjne lub czujniki pływakowe – do wykrywania wycieków (pralka, piwnica, serwerownia). Metody zestawia tabela 2. **WM**

Projekty pokrewne i literatura

Projekty w polskiej prasie

- **Elektronika Praktyczna – Agri-Stick – urządzenie do monitorowania parametrów gleby.** Omawia sondę pojemnościową (odporną na korozję) z wyjściem napięciowym oraz pomiar wilgotności i temperatury powietrza czujnikiem DHT22 – dobre porównanie z rezystancyjnym rozwiązaniem z części A.
<https://ep.com.pl/projekty/projekty-soft/12911-agri-stick-urzadzenie-domonitorowania-parametrow-gleby-wrolnictwie>
- **Elektronika Praktyczna – Internet Rzeczy w pomiarach środowiskowych (6): zestaw Enviro Grow.** Pojemnościowe czujniki wilgoci gleby z wyjściem cyfrowym (modulacja częstotliwości) oraz pompki i czujnik parametrów powietrza – nowoczesne, sieciowe podejście do podlewania. <https://ep.com.pl/projekty/moduly-w-aplikacjach/16161-internet-rzeczy-w-pomiarach-srodowiskowych-6-zestaw-do-kontrolowania-wilgotnosc-gleby-enviro-grow-firmy-pimoroni>

Poradniki i konstrukcje

- **elektroweb.pl – Czujnik wilgotności gleby z Arduino** – poradnik krok po kroku. praktyczne podłączenie taniego modułu pojemnościowego (HW-390) do Arduino/ESP: zasilanie 3,3 V, dopasowanie odniesienia ADC (AREF) i kalibracja odczytu.
<https://elektroweb.pl/projekty/czujnik-wilgotnosc-gleby-z-arduino/>
- **Pino-Tech (PL) – SoilWatch 10 – pojemnościowy czujnik wilgotności gleby.** Polski czujnik z własnym regulatorem napięcia (niezależnia sygnał od wahań zasilania) i częstotliwością pracy dobraną tak, by ograniczyć wpływ zasolenia; opisany w cyklu „Internet Rzeczy w pomiarach środowiskowych” w EP. <https://ep.com.pl/projekty/moduly-w-aplikacjach/16161-internet-rzeczy-w-pomiarach-srodowiskowych-6-zestaw-do-kontrolowania-wilgotnosc-gleby-enviro-grow-firmy-pimoroni>

Moduł pojemnościowy „soil moisture v2.0” (standard DIY)

- **Moduł „soil moisture v1.2/v2.0”** – pojemnościowy czujnik wilgotności gleby. koplanarne ścieżki tworzą kondensator, a generator na układzie 555 (~1,5 MHz) z prostownikiem i filtrem RC daje na wyjściu napięcie analogowe (typowo ~1,2...2,5 V między „w wodzie” a „w powietrzu”). Brak kontaktu galwanicznego zapewnia odporność na korozję i 2...3× dłuższą żywotność niż sonda rezystancyjna. Warto sprawdzić dwie rzeczy w tanich egzemplarzach: układ czasowy powinien być CMOS-owy TLC555 (wersje z NE555 bywają niestabilne przy 3,3 V), a stabilizator 3,3 V (662k) bywa zastąpiony rezystorem – wtedy napięcie wyjściowe zależy od zasilania i wymaga uśredniania odczytu. Analiza i wskazówki: <https://www.makerguides.com/capacitive-soil-moisture-sensor-with-arduino/>; <https://thecavepearlproject.org/2020/10/27/hacking-a-capacitive-soil-moisture-sensor-for-frequency-output/>

Czujniki wilgotności powietrza (noty katalogowe)

- **Sensirion – SHT4x – cyfrowy czujnik wilgotności i temperatury.** Skalibrowany higrometr I²C do pomiaru wilgotności powietrza. <https://sensirion.com/products/catalog/SHT40>
- **Bosch Sensortec – BME280 – czujnik wilgotności, temperatury i ciśnienia.** Popularny układ środowiskowy do stacji pogodowych i monitoringu wnętrz. <https://www.bosch-sensortec.com/products/environmental-sensors/humidity-sensors-bme280/>

Gotowe czujniki pojemnościowe

- **Adafruit – STEMMa Soil Sensor – I²C Capacitive Moisture Sensor.** Pomiar pojemnościowy realizowany wbudowanym układem pomiaru dotyku mikrokontrolera ATSAM10 (odczyt od ~200 „sucho” do ~2000 „mokra”), interfejs I²C, dodatkowo temperatura; brak odsonionego metalu, więc bez korozji i prądów statycznych w glebie. <https://learn.adafruit.com/adafruit-stemma-soil-sensor-i2c-capacitive-moisture-sensor>
- **SparkFun – Qwiic Soil Moisture Sensor (CY8CMBR3102 CapSense).** Pojemnościowa płytka z kontrolerem CapSense i progową diodą stanu, z podziałką głębokości i biblioteką kalibracji (Arduino/MicroPython). Podobną, tanią sondą analogową jest DFRobot Gravity (wyjście 1,2...2,5 V, odporna na korozję). https://docs.sparkfun.com/SparkFun_Capacitive_Soil_Moisture_Sensor_CY8CMBR3102/introduction

Sondy profesjonalne i pomiar gruntu (TDR/FDR)

- **Vegetronix – VH400 – sonda wilgotności gleby (technika linii transmisyjnej/TDR).** Pomiar objętościowej zawartości wody (VWC 0...50%), wyjście napięciowe 0...3 V, odporność na zasolenie i brak odsonionego metalu (nie koroduje). Punktem odniesienia w badaniach są też sondy FDR METER Group/Decagon EC-5. <https://www.vegetronix.com/Products/VH400/>

arXiv – smol: Sensing Soil Moisture using LoRa. przykład użycia miernika FieldScout TDR-300 jako odniesienia (TDR – pomiar czasu przelotu fali elektromagnetycznej w gruncie) przy kalibracji sieci czujników. <https://arxiv.org/pdf/2110.01501>

Wykrywanie zalania i wycieków

Ta sama zasada progowa, co w części A, skaluje się od detekcji punktowej po czujniki liniowe (kable), które wykrywają ciecz na całej swojej długości.

- **IndustrialMonitorDirect – Water Leak Sensing Cable** – zasada działania i projekt układu. liniowy czujnik wykrywający wodę w dowolnym miejscu na długości kabla: topologia mostka Wheatstone’a, próg na przetworniku ADC oraz lokalizacja wycieku metodą reflektometrii (TDR); impedancja falowa kabla $Z_0 \approx 100 \Omega$. Dobre rozwinięcie idei prostego detektora obecności wody. <https://industrialmonitordirect.com/blogs/knowledgebase/water-leak-sensing-cable-detection-principle-and-circuit-design>
- **HW group – Sensor WLD Relay – czujnik zalania z wyjściem przekaźnikowym.** Przemysłowy detektor z kablem czujnikowym (jedna strefa nawet do 185 m), sygnalizujący stany OK/woda/przerwa stykiem przekaźnika (NO/NC) – komercyjny odpowiednik wyjścia przekaźnikowego z części A (por. też rozwiązania NTI ENVIROMUX i BAPI). <https://www.hw-group.com/sensor/sensor-wld-relay-1w-uni>

Czujniki wilgotności i cieczy w ofercie kitów AVT

Projekt przykładowy nie jest zestawem AVT, ale tematyka wilgotności i wykrywania cieczy jest w ofercie AVT obecna. Poniżej zestawiamy najbliższe tematycznie pozycje ze sklepu AVT (sklep.avt.pl).

- **AVT764 Czujnik wilgoci i zalania z alarmem dźwiękowym.** Prosty sygnalizator reagujący na wodę, deszcz lub zwarcie elektrod; brzęczyk piezo, zasilanie z baterii CR2032, pobór w czuwaniu poniżej 1 μ A. Elektrody w formie grzebienia – bliski krewny sondy z części A.
- **Moduł czujnika wilgotności gleby** (komponent). gotowa sonda do automatyki nawadniania roślin doniczkowych; zasilanie 3,3...5 V.
- Od dwóch kawałków metalu w doniczce po pojemnościowe sondy i cyfrowe higrometry – pomiar wilgotności pozostaje projektem, w którym prosty układ progowy potrafi realnie oszczędzić wodę, roślinę albo zalaną podłogę.

Bibliografia

- [1] B. Landoni, „Soil Moisture Sensor”, Open-Electronics.org – <https://www.open-electronics.org/soil-moisture-sensor/>
- [2] Elektronika Praktyczna, „Agri-Stick – urządzenie do monitorowania parametrów gleby” – <https://ep.com.pl/projekty/projekty-soft/12911-agri-stick-urzadzenie-domonitorowania-parametrow-gleby-wrolnictwie>
- [3] Elektronika Praktyczna, „Internet Rzeczy w pomiarach środowiskowych (6): Enviro Grow/SoilWatch 10” – <https://ep.com.pl/projekty/moduly-w-aplikacjach/16161-internet-rzeczy-w-pomiarach-srodowiskowych-6-zestaw-do-kontrolowania-wilgotnosci-gleby-enviro-grow-firmy-pimoroni>
- [4] elektroweb.pl, „Czujnik wilgotności gleby z Arduino – poradnik” – <https://elektroweb.pl/projekty/czujnik-wilgotnosci-gleby-z-arduino/>
- [5] Makerguides, „How to Use a Capacitive Soil Moisture Sensor with Arduino” – <https://www.makerguides.com/capacitive-soil-moisture-sensor-with-arduino/>
- [6] The Cave Pearl Project, „Hacking a Capacitive Soil Moisture Sensor for Frequency Output” – <https://thecavepearlproject.org/2020/10/27/hacking-a-capacitive-soil-moisture-sensor-for-frequency-output/>
- [7] Adafruit, „STEMMA Soil Sensor – I²C Capacitive Moisture Sensor” – <https://learn.adafruit.com/adafruit-stemma-soil-sensor-i2c-capacitive-moisture-sensor>
- [8] SparkFun, „Qwiic Soil Moisture Sensor (CY8CMBR3102)” – https://docs.sparkfun.com/SparkFun_Capacitive_Soil_Moisture_Sensor_CY8CMBR3102/introduction
- [9] Vegetronix, „VH400 Soil Moisture Sensor” – <https://www.vegetronix.com/Products/VH400/>
- [10] „smol: Sensing Soil Moisture using LoRa”, arXiv – <https://arxiv.org/pdf/2110.01501>
- [11] Sensirion, „SHT4x – czujnik wilgotności i temperatury” – <https://sensirion.com/products/catalog/SHT40>
- [12] Bosch Sensortec, „BME280 – czujnik wilgotności, temperatury i ciśnienia” – <https://www.bosch-sensortec.com/products/environmental-sensors/humidity-sensors-bme280/>
- [13] IndustrialMonitorDirect, „Water Leak Sensing Cable – Detection Principle and Circuit Design” – <https://industrialmonitordirect.com/blogs/knowledgebase/water-leak-sensing-cable-detection-principle-and-circuit-design>
- [14] HW group, „Sensor WLD Relay 1W-UNI – Water Leak Detector” – <https://www.hw-group.com/sensor/sensor-wld-relay-1w-uni>



Druk 3D w służbie elektroniki (11)

Po dłuższej przerwie wracamy do tematyki druku 3D z kolejnym praktycznym projektem. Tym razem skupimy się na konstrukcji, która musi być wytrzymała mechanicznie, a jednocześnie możliwie najlżejsza. W tym celu zostaną zastosowane różne materiały, w tym jeden zaawansowany filament techniczny, oraz metoda wyżarzania w osłonie gipsowej. Co będziemy drukować? Szkielet oraz osłony aerodynamiczne dużego drona FPV.

Zreguły metodą druku 3D wykonuje się drony klasy mikro oraz konstrukcje o masie startowej poniżej 500 g. W przypadku cięższych platform bezzałogowych najczęściej wykorzystuje się ramy wycinane na frezarkach CNC z płyt laminatu z włókna węglowego (karbonu). Laminaty te charakteryzują się niską masą, wysoką sztywnością i znakomitą wytrzymałością mechaniczną, choć w razie twardego lądowania lub zderzenia ulegają pęknięciom bądź rozwarstwieniu (delaminacji). Włókno węglowe wykazuje wybitną wytrzymałość na rozciąganie wzdłuż włókien, natomiast odporność na pozostałe kierunki naprężeń zależy bezpośrednio od parametrów żywicy spajającej kompozyt. O ile gotowe, seryjne ramy FPV są relatywnie tanie dzięki masowej produkcji (fotografia przedstawia ramę 7" dostępną za mniej niż 100 zł), o tyle wykonanie prototypowej konstrukcji na zamówienie bywa kosztowne. W takich warunkach druk 3D stanowi atrakcyjną alternatywę, należy jednak uwzględnić anizotropię właściwości mechanicznych wydruków w porównaniu z litym laminatem. Kluczowy jest odpowiedni dobór materiału oraz optymalizacja geometrii elementów nośnych. Spotyka się również podejście hybrydowe, w którym fabryczna płyta węglowa

stanowi rdzeń nośny, a elementy drukowane nadają całości docelową formę i aerodynamikę. Wymaga to jednak dysponowania precyzyjnym modelem CAD ramy bazowej i narzuca ograniczenia geometryczne, np. w kwestii rozstawu osi silników.

Założenia techniczne projektu oraz dobór komponentów

Typowy szkielet drona wyścigowego sprawdza się znakomicie w lekkich konstrukcjach akrobacyjnych. Silniki i stos elektroniki (*stack*) przykręca się bezpośrednio do ramy, moduły łączności (odbiornik aparatury, nadajnik wideo oraz moduł GNSS) mocuje się taśmą dwustronną, a pakiet zasilający podwiesza opaską rzepową. Konfiguracja ta pozwala na kilkuminutowy, dynamiczny lot. Celem niniejszego projektu jest jednak zbudowanie platformy długodystansowej (typu long range), wyposażonej w pojemny akumulator oraz dodatkowe moduły telemetryczne, w tym miniaturowy LIDAR do wykrywania przeszkód.

Projekt zakłada wykorzystanie rozwiązań aerodynamicznych w celu poprawy sprawności napędu i wydłużenia

czasu lotu. Wirniki zostaną umieszczone w tunelach pierścieniowych (wentylatory kanałowe, *ducted fans*), co w teorii pozwala zwiększyć generowany ciąg statyczny nawet o 30%. Eksperymentalnie zaimplementowane zostaną także profile wykorzystujące efekt Coandy. Zjawisko to polega na przyleganiu strumienia płynu (powietrza) do zakrzywionej powierzchni opływanej. Odpowiednie ukształtowanie profilu kierownicy strugi skupia strumień zaśmigłowy, dając dodatkowy przyrost ciągu rzędu 5...10% poprzez zaciąganie mas powietrza z otoczenia strugi. Zamknięcie śmigieł w tunelach zwiększa jednak powierzchnię boczną drona, co przekłada się na większą podatność na podmuchy wiatru oraz zmniejszenie zwrotności. W przypadku projektowanej platformy wysokość prędkość przelotowa i zdolności akrobatyczne są jednak parametrami drugorzędnymi.

Do napędu drona posłużą cztery bezszczotkowe silniki prądu stałego (BLDC) typu outrunner – model ECO II 2807 marki Emax. Silniki te przystosowane są do zasilania z pakietu 6S i charakteryzują się współczynnikiem

$$K_v = 1300 \text{ obr.} / (\text{min} \cdot V)$$

Będą one współpracować z trójęłatowymi śmigłami o wymiarach 7×3,5×3 (średnica 7 cali, skok 3,5 cala). Taki skok łopat generuje wysoki ciąg w konfiguracji tunelowej, co wymusza zastosowanie jednostek o niższej prędkości obrotowej i wyższym momencie obrotowym. Do sterowania wybrano kontroler lotu SpeedyBee F405 V4 zintegrowany z czterokanałowym regulatorem ESC o obciążalności ciągłej 55 A na kanał. Pojedynczy silnik pod pełnym obciążeniem może pobierać do 52 A, co przy napięciu pakietu 6S (25,2 V) daje moc rzędu 1310 W na kanał. Do zasilania układu posłuży pakiet litowo-jonowy 6S3P o pojemności 12 000 mAh i wydajności prądowej 10 C. Przy nominalnej wydajności 120 A akumulator nie pokryłby pełnego, jednoczesnego zapotrzebowania czterech silników pracujących na 100% mocy. Zastosowany zespół napędowy generuje jednak sumaryczny ciąg przekraczający 8 kg (bez uwzględnienia zysku z kanałów), a z dyszami nawet 11 kg. Przy szacowanej masie startowej drona poniżej 3 kg stanowi to znaczny naddatek mocy. Z tego względu w oprogramowaniu układowym Betaflight wprowadzony zostanie programowy limit wysterowania przepustnicy na poziomie 75...80%. Pozwoli to uzyskać stosunek ciągu do masy (TWR, *Thrust to Weight Ratio*) na poziomie 3,0...3,2:1, co stanowi wartość w zupełności wystarczającą do stabilnych, płynnych lotów z kamerą.

Układ pozycjonowania i stabilizacji przestrzennej bazyje na trzech niezależnych modułach sensorów. Pierwszym z nich jest miniaturowy LIDAR skanujący YDLIDAR X2L. Charakteryzuje się on zakresem pomiarowym 0,12...8 m, niepewnością pomiaru ±2 cm oraz rozdzielczością kątową



Fotografia 1. Przemysłowy szkielet ramy drona FPV 7" (Mark 4) wykonany z frezowanego laminatu węglowego

ok. 0,84°. Ze względu na ograniczone zasoby obliczeniowe głównego kontrolera F405, do akwizycji i wstępnej filtracji danych z LIDAR-u oddelegowano dedykowany mikrokontroler dsPIC. Jego zadaniem jest wyliczanie wektorów unikania kolizji i przekazywanie ich w formie uproszczonych ramek korygujących do magistrali kontrolera lotu. Do precyzyjnego utrzymywania pozycji w zawisie przewidziano optyczny czujnik przepływu (Optical Flow) ThoneFlow 3901U, współpracujący z impulsowym dalmierzem laserowym mierzącym rzeczywistą wysokość nad podłożem (AGL). Integrację danych telemetrycznych z czujnika przepływu, dalmierza oraz LIDAR-u przeprowadza układ dsPIC, odciażając główny procesor nawigacyjny. Trzecim elementem nawigacyjnym jest moduł GNSS BN-880 zintegrowany z magnetometrem cyfrowym (kompasem), oparty na odbiorniku u-blox M8030-KT. Układ ten zapewnia równoległą obsługę dwóch systemów pozycjonowania satelitarne (GPS + GLONASS lub GPS + BeiDou), osiągając horyzontalną dokładność wyznaczenia pozycji rzędu 2 m.

Tor wizyjny oraz aparatura RC bazują na sprawdzonych modułach. Zastosowano kamerę analogową Caddx Ratel 2 V2 (format obrazu 4:3/16:9, system PAL/NTSC) ze względu na zminimalizowane opóźnienia przesyłu oraz korzystny bilans cenowy. Sygnał wideo trafia na płytkę kontrolera lotu, gdzie układ scalony AT7456E (zamiennik układu MAX7456) nakłada warstwę OSD z danymi telemetrycznymi (napięcie baterii, prąd pobierany, współrzędne, wysokość, prędkość). Zmiksowany sygnał jest transmitowany przez nadajnik VTX SpeedyBee TX800 pracujący w paśmie 5,8 GHz z mocą do 800 mW. Do odbioru sygnałów sterujących służy miniaturowy odbiornik RadioMaster RP1 V2 pracujący w paśmie 2,4 GHz z oprogramowaniem ExpressLRS (protokół CRSF przez UART). Zbudowano go w oparciu o mikrokontroler ESP8285 (rdzeń Tensilica L106, taktowanie 80 MHz) oraz układ transceivera Semtech SX1280 (modulacja LoRa), co gwarantuje daleki zasięg łącząca sterującego i wysoką odporność na zakłócenia.

Dobór filamentów

Konstrukcja nośna drona musi łączyć wysoką sztywność geometryczną (niezbędną do uniknięcia drgań rezonansowych rejestrowanych przez żyroskopy) z odpornością na udary mechaniczne. W klasycznych konstrukcjach laminat węglowo-epoksydowy charakteryzuje się modułem Younga E w granicach 50...60 GPa oraz wytrzymałością na zginanie σ na poziomie 400...500 MPa. Wadą kompozytu węglowego jest jednak stosunkowo niska odporność na obciążenia ścinające żywicę oraz kruchość przy uderzeniach punktowych, co w bezzałogowcach prowadzi do rozwarstwiania ramion silników przy twardych zderzeniach.

W projekcie poddano szczegółowej analizie dwa tworzywa termoplastyczne: PCTG oraz ASA-GF (kopolimer akrylonitryl-styren-akrylan domieszkowany ciętym włóknem szklanym). Czysty PCTG charakteryzuje się znakomitą udatnością (próbki bez karbu nie ulegają pęknięciu w próbie udarowej), jednak jego moduł sprężystości wzdłużnej wynosi zaledwie 2,0...2,2 GPa przy wytrzymałości na zginanie rzędu 75...85 MPa. Wykonanie długich ramion silników wyłącznie z PCTG skutkowałoby zbyt dużą podatnością giętą i generowaniem wibracji o niskiej częstotliwości. Z kolei filament ASA-GF oferuje moduł Younga na poziomie 3,5...4,5 GPa oraz wytrzymałość na zginanie 60...80 MPa. Dodatek włókna szklanego znacząco ogranicza skurcz przetwórczy i zwiększa stabilność wymiarową detali. Gęstość kompozytu węglowego wynosi 1,5...1,6 g/cm³, gęstość PCTG to 1,20 g/cm³, a stopnia ASA-GF mieści się w zakresie 1,11...1,15 g/cm³.

Główne elementy nośne szkieletu drona zostaną wydrukowane z filamentu ASA-GF. W celu usunięcia naprężeń termicznych, zwiększenia krystaliczności i podniesienia adhezji międzywarstwowej elementy te zostaną zasypane w drobnym gipsie formierskim i poddane wyżarzaniu (annealingowi) w temperaturze 125...130°C. Zabieg ten pozwala podnieść graniczną wytrzymałość na zginanie o 20...35% przy zachowaniu stabilności wymiarowej zabezpieczonej formą gipsową. Z kolei elementy montażowe pakietu, osłony elektroniki oraz tunele aerodynamiczne zostaną wykonane z udurowionego PCTG bez obróbki cieplnej. Rozważane zastosowanie elastomeru TPU porzucono ze względu na zbyt wysoką elastyczność, która w tunelach śmigieł doprowadzałaby do ocierania łopatek o ścianki pod wpływem podciśnienia.

Podsumowanie

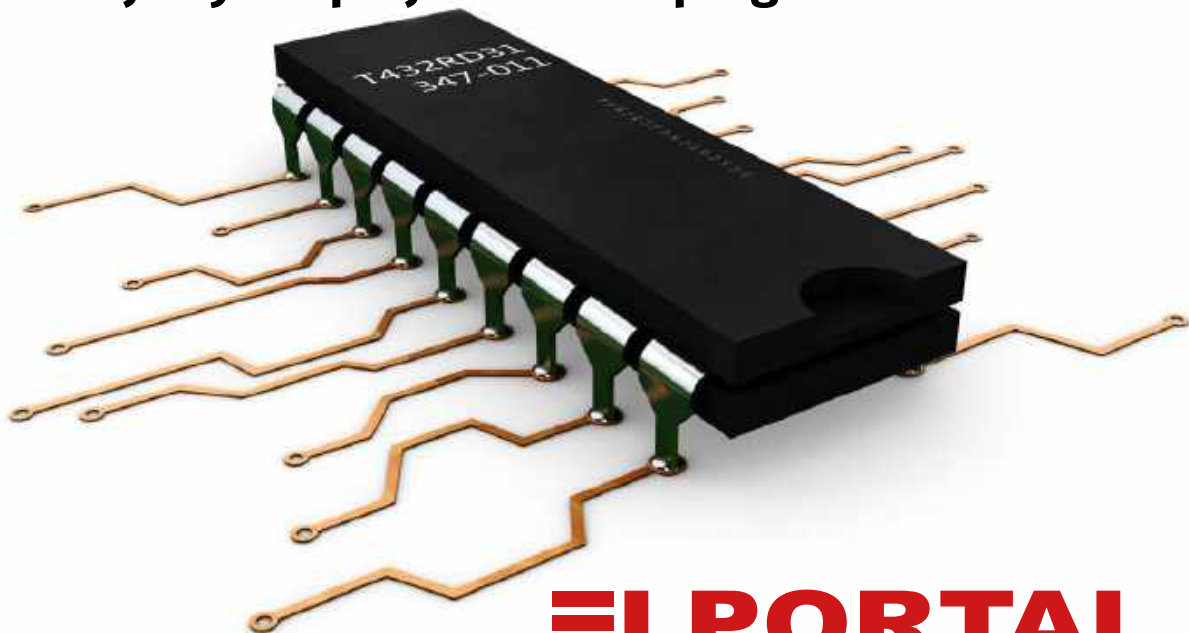
Na etapie koncepcyjnym analizowano także inne materiały: PCTG-GF (trudno dostępny handlowo), PLA-CF (niedostateczna odporność termiczna podczas ekspozycji na słońcu i ciepło silników) oraz wysokotemperaturowe kompozyty PA-CF i PC-CF (wymagające zaawansowanych komór grzejnych i generujące wysoki koszt materiałowy). Odrzucono również pomysł zastosowania profili aluminiowych V-Slot ze względu na niekorzystny bilans masowy.

W kolejnej części cyklu przedstawiony zostanie kompletny projekt ramy CAD, analiza masowa poszczególnych węzłów konstrukcyjnych oraz szczegółowe parametry procesu druku i wyżarzania w piecu laboratoryjnym.

Paweł Kowalczyk

REKLAMA

Publikujemy dla projektantów i programistów elektroniki



ELPORTAL.pl



W ofercie AVT*

AVT6102

Najważniejsze parametry:

- załączanie jednego urządzenia po wykryciu poboru prądu przez drugie
- wyłączanie sterowanego urządzenia z regulowanym opóźnieniem 0...10 s
- regulowany próg natężenia prądu wywołującego załączenie
- zasilanie z sieci 115/230 V
- maksymalny pobór prądu przez urządzenie sterowane: 16 A
- maksymalny łączny pobór prądu przez urządzenie nadrzędne i sterowane: 20 A
- minimalny pobór prądu przez urządzenie nadrzędne: 0,4 A

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

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

Projekty pokrewne na stronie www.ep.com.pl

- (aktywne linki do artykułów):
- Mikrokłaskacz – mikroprocesorowy włącznik akustyczny
 - Zdalny włącznik radiowy
 - Podwójny włącznik dotykowy
 - Automatyczny włącznik zmierzchowy
 - Zbliżeniowy włącznik refleksyjny
 - Włącznik wielowejściowy

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

Włącznik pomocniczy

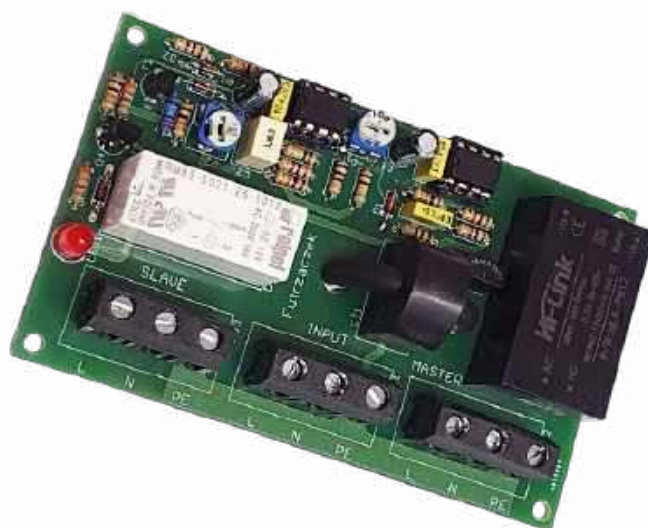
Często zdarza się, że jedno urządzenie elektryczne musi pracować dokładnie wtedy, kiedy działa inne – na przykład odkurzacz przy włączonej wyrzynarce lub wyciąg oparów podczas pracy lutownicy. Prezentowany układ pozwala zsynchronizować ich załączanie oraz zapewnia wyłączenie urządzenia sterowanego z regulowanym opóźnieniem.

Wiercenie otworów w ścianie remontowanego pomieszczenia to żaden problem, wszak kurz i tak jest wszędzie. Gorzej, kiedy trzeba zawiesić półkę w wykończonym już salonie. Jednoczesne wiercenie oraz włączanie i wyłączanie odkurzacza zbierającego pył bywa kłopotliwe. Podobnie wygląda sytuacja z lutowaniem i uruchamianiem wentylatora wyciągającego opary. Przykładów można podać wiele, a łączy je jedno: urządzenie pomocnicze powinno uruchamiać się automatycznie wraz ze sprzętem głównym i wyłączać, gdy ten przestaje pracować.

Układ realizuje tę funkcję we współpracy z odbiornikami zasilanymi jednofazowym napięciem przemiennym z sieci elektroenergetycznej. Ponieważ urządzenie nadrzędne może pobierać niewielki prąd w trybie czuwania przez cały czas, wprowadzono regulację progu czułości – minimalnego natężenia prądu, powyżej którego następuje reakcja. W wielu zastosowaniach celowe jest także podtrzymanie pracy odbiornika sterowanego (np. odkurzacza) przez pewien czas po wyłączeniu narzędzia głównego, co w tym rozwiązaniu jest bardzo proste do uzyskania.

Budowa

Schemat ideowy układu przedstawiono na **rysunku 1**. Do detekcji prądu pobieranego przez urządzenie nadrzędne służy przekładnik prądowy CT1 obciążony rezystorem R1. Wzmacniacz operacyjny US1A prostuje jednopółkrowo

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

- R1: 100 Ω
- R2, R3, R8, R12: 1 kΩ
- R4, R7, R9, R11, R15: 10 kΩ
- R5, R13: 100 kΩ
- R6: 220 Ω
- R10: 22 Ω
- R14, R16: 3 kΩ
- P1: 2 kΩ (potencjometr montażowy leżący)
- P2: 2 MΩ (potencjometr montażowy leżący)

Kondensatory:

- C1: 10 nF/63 V (raster 5 mm, MKT)
- C2: 2,2 μF/63 V (raster 5 mm, MKT)
- C3, C5: 100 nF/63 V (raster 5 mm, MKT)
- C4, C6: 22 μF/25 V (elektrolityczny, raster 2,5 mm)

Półprzewodniki:

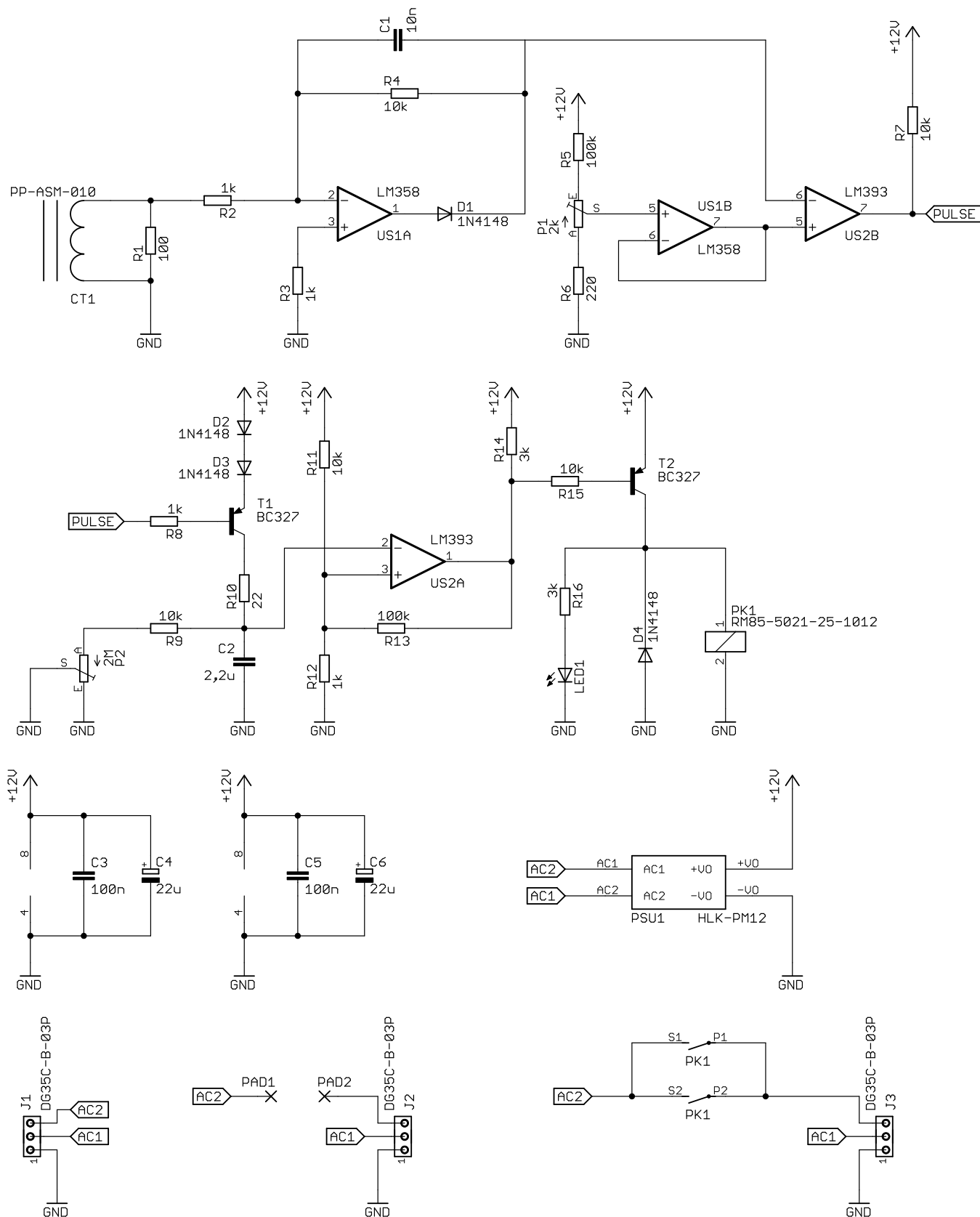
- D1...D4: 1N4148
- LED1: LED 5 mm czerwona (np. LED F5 R)
- T1, T2: BC327
- US1: LM358 (DIP8)
- US2: LM393 (DIP8)

Pozostałe:

- CT1: PP-ASM-010
- J1...J3: DG35C-B-03P
- PK1: RM85-5021-25-1012
- PSU1: HLK-PM12
- Odcinek izolowanego przewodu miedzianego 2,5 mm²
- Podstawki DIP8 – 2 szt.

i wzmacnia sygnał z przekładnika ze wzmacnieniem 10 V/V, ustalonym stosunkiem rezystorów R_5/R_4 (dla ujemnych połówek napięcia; dodatnie są przepuszczane bez wzmacnienia). Dioda D1 ułatwia uzyskanie potencjału 0 V na wyjściu wzmacniacza operacyjnego. Ponieważ jest on zasilany niesymetrycznie, jego stopień wyjściowy w stanie

nasylenia nie osiąga napięcia niższego niż kilkanaście miliwoltów. Włączenie szeregowo diody D1 sprawia, że wyjście US1A musi osiągnąć potencjał co najmniej kilkaset miliwoltów, aby podnieść napięcie na katodzie powyżej zera. Kondensator C1 ogranicza pasmo przenoszenia, filtrując zakłócenia wielkiej częstotliwości.



Rysunek 1. Schemat ideowy włącznika pomocniczego

Za detekcję przekroczenia zadanego progu odpowiada komparator US2B (popularny układ LM393). Porównuje on wzmacnione napięcie z US1A z napięciem odniesienia pobieranym ze ślizgacza potencjometru P1. Wtórnik napięciowy US1B nie bierze udziału w torze sygnałowym, lecz jego wejścia zostały spolaryzowane, aby uniknąć pozostawienia niepodłączonych wyprowadzeń.

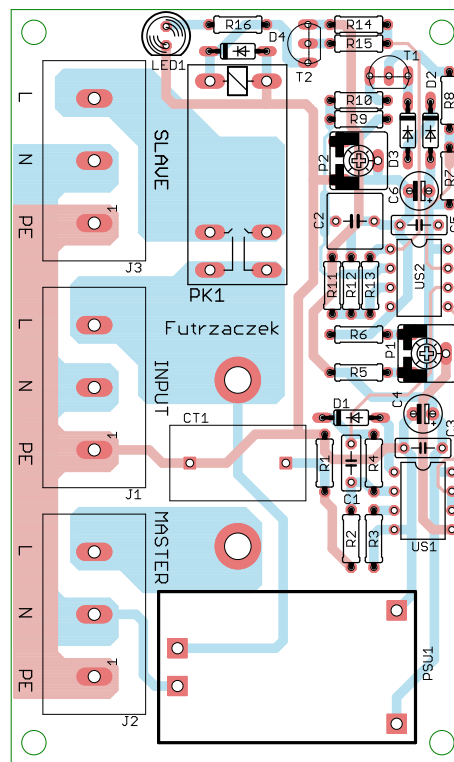
Gdy prąd płynący przez przewód przewleczony przez przekładnik (między punktami lutowniczymi PAD1 i PAD2) przekroczy próg ustawiony potencjometrem P1, na wyjściu komparatora US2B pojawia się stan niski. Powoduje to otwarcie tranzystora T1, z którego bazy wypływa prąd ograniczany przez rezystor R8 do wartości ok. 10 mA. Tak znaczny prąd bazy pozwala na szybkie ładowanie kondensatora C2 dużym prądem kolektora. Rezystor R10 ogranicza ten prąd do wartości bezpiecznej dla struktury tranzystora. Diody D2 i D3 włączone szeregowo w obwód emitera obniżają maksymalny potencjał ładowania kondensatora C2 o sumaryczną wartość napięcia przewodzenia (ok. 1,5 V poniżej szyny zasilania), co spełnia wymogi wejściowe komparatora LM393. Rezystor R7 polaryzuje bazę T1 i utrzymuje go w stanie zatkania, gdy wyjście typu otwarty kolektor komparatora US2B pozostaje nieaktywne.

Otwarcie tranzystora T1 – nawet na krótki czas – ładuje kondensator C2. Rezystancja rezystora R9 połączonego szeregowo z potencjometrem P2 odpowiada za powolne rozładowanie C2, a stała czasowa tego obwodu wyznacza czas opóźnienia wyłączenia odbiornika sterowanego: im większa rezystancja, tym dłuższy czas podtrzymania zasilania.

Napięcie z kondensatora C2 trafia na wejście komparatora US2A. Napięcie odniesienia jest formowane przez dzielnik rezystorowy R11 i R12 zasilany z głównej szyny 12 V. Rezystor R13 wprowadza dodatnie sprzężenie zwrotne, tworząc histerezę i eliminując oscylacje wyjścia przy wolnych zmianach napięcia na C2. W komparatorze US2B histereza nie była konieczna, gdyż pożądane jest doładowywanie C2 nawet pojedynczymi, wąskimi impulsami przy pracy na granicy progu detekcji.

Gdy napięcie na C2 przekracza próg komparacji, na wyjściu US2A pojawia się stan niski, co wprowadza tranzystor T2 w stan nasycenia (prąd bazy jest ograniczany rezystorem R15 do ok. 1 mA). Rezystor R14 pełni kilka funkcji: podciąga wyjście komparatora do szyny zasilania, zapewnia prąd polaryzacji stopnia wyjściowego komparatora, utrzymuje T2 w stanie zatkania przy braku wystereowania oraz polaryzuje rezystor histerezy R13 połączony z dzielnikiem R11/R12. Wartość R14 dobrano tak, aby spadek napięcia przy zatkanym komparatorze nie przekraczał 0,7 V, co gwarantuje pewne wyłączenie tranzystora T2.

Nasycony tranzystor T2 załącza przełącznik PK1, włączając obwód urządzenia sterowanego. Równolegle



Rysunek 2. Schemat montażowy i wzór ścieżek płytki drukowanej

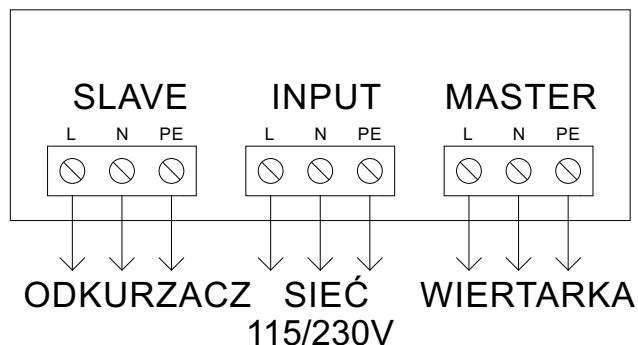
z cewką przekładnika zasilana jest dioda LED1 sygnalizująca stan pracy. Dioda D4 zabezpiecza tranzystor T2 przed przepięciami indukcyjnymi generowanymi podczas wyłączania cewki przekładnika.

Część pomiarowo-sterująca jest galwanicznie odizolowana od sieci elektroenergetycznej za pośrednictwem przekładnika prądowego CT1 oraz miniaturowego zasilacza impulsowego PSU1 o napięciu wyjściowym 12 V. Zacisk przewodu ochronnego (PE) połączono z masą układu (GND).

Montaż i uruchomienie

Układ zmontowano na dwustronnej płytce drukowanej o wymiarach 100 mm × 60 mm. Rozmieszczenie elementów oraz mozaikę ścieżek przedstawiono na **rysunku 2**. W odległości 3 mm od krawędzi płytki przewidziano cztery otwory montażowe o średnicy 3,2 mm.

Montaż warto rozpocząć od elementów najniższych – rezystorów i diod półprzewodnikowych. Przez otwór przekładnika CT1 należy przewlec odizolowany na końcach odcinek przewodu miedzianego, wlutować jego końce w punkty lutownicze PAD1 i PAD2 przed i za przekładnikiem, a następnie przylutować sam przekładnik. Pod układy scalone US1 i US2 zaleca się zastosowanie podstawek precyzyjnych. Szerokie ścieżki prądowe, pozbawione maski lutowniczej, należy obficie pobielić spoiwem lutowniczym w celu zwiększenia ich przekroju i obciążalności prądowej. Zamiast złączy śrubowych DG35C-B-03P (J1...J3) można zastosować standardowe złącza w rastrze 7,5 mm (wymaga to delikatnego rozgięcia wyprowadzeń), pamiętając jednak o ich niższej



Rysunek 3. Schemat aplikacji i podłączenia urządzeń zewnętrznych

obciążalności prądowej. Zmontowany prototyp przedstawiono na fotografii otwierającej.

Prawidłowo zmontowany moduł działa od razu po włączeniu. Do złącza J1 (INPUT) doprowadza się napięcie sieciowe. Zasilanie to trafia bezpośrednio na zaciski złącza J2 (MASTER), a układ stale monitoruje pobierany prąd i na tej podstawie steruje zasilaniem złącza J3 (SLAVE). Sposób podłączenia odbiorników zilustrowano na rysunku 3.

Do regulacji parametrów pracy służą dwa potencjometry montażowe. Potencjometr P1 ustala próg czułości detekcji prądu: przy skrajnym lewym położeniu ślizgacza układ reaguje na prąd o wartości skutecznej od ok. 0,4 A; obrót w prawo zmniejsza czułość (podnosząc próg do ok. 3 A), co zapobiega fałszywym załączeniom przy odbiornikach pobierających prąd w trybie czuwania. Potencjometr P2 reguluje czas podtrzymania zasilania wyjścia SLAVE po wyłączeniu urządzenia MASTER: skrajne prawe położenie oznacza natychmiastowe rozłączenie styków, natomiast obrót w lewo wydłuża ten czas do ok. 10 s.

Maksymalny ciągły prąd pobierany przez odbiornik sterowany (SLAVE) wynosi 16 A, co wynika z dopuszczalnej obciążalności styków przekaźnika PK1. Całkowity prąd przepływający przez złącze zasilające J1 nie może przekraczać 20 A przy zastosowaniu złącz DG35C-B-03P i odpowiednim pogrubieniu ścieżek zasilających spoiwem.

Michał Kurzela

świat radio
Magazyn wszystkich użytkowników eteru
KROTKOFALARSTWO CB RADIOTECHNIKA

przejrzyj i kupisz na stronie
www.ulubionykiosk.pl

ŁOŚ oraz inne spotkania i zjazdy krotkofalarów
świat radio
Magazyn wszystkich użytkowników eteru
KROTKOFALARSTWO CB RADIOTECHNIKA
9-10-20
Wzmacniacz mocy ICOM IC-PW2

Obudowa to część układu

Jak uniknąć problemów, których nie widać na schemacie?

Za mało miejsca na przewody, kłopotliwy montaż, utrudniona wymiana modułu albo nadmierna temperatura komponentów często ujawniają się dopiero wtedy, gdy konstrukcja urządzenia jest już zaawansowana. Ich źródłem nie musi być jednak błąd w schemacie lub projekcie PCB. Nierzadko problem zaczyna się od potraktowania obudowy jako elementu dobieranego dopiero po zaprojektowaniu elektroniki. Tymczasem jej geometria, sposób montażu, technologia przyłączeniowa i właściwości materiału wpływają na całe urządzenie.

Zewnętrzny wymiar obudowy nie mówi, ile miejsca zajmie urządzenie

Dobór obudowy na podstawie długości, szerokości i wysokości to dopiero początek. W rzeczywistej aplikacji trzeba uwzględnić także przestrzeń potrzebną na podłączenia wejść i wyjść czy zasilania, przewody i ich promienie gięcia oraz dostęp do punktów przyłączeniowych. Znaczenie ma również umiejscowienie urządzenia na szynie DIN i możliwość jego późniejszego demontażu.

W szafie sterowniczej obudowa funkcjonuje w otoczeniu kanałów kablowych, innych urządzeń i elementów konstrukcyjnych. Kierunek wyprowadzenia przewodów może więc przesądzić o wymaganej głębokości szafy, odstępach od kanału kablowego i liczbie urządzeń mieszczących się na jednej szynie. Nawet kompaktowy moduł może w praktyce potrzebować znacznie więcej miejsca, jeżeli przyłącza wymagają prowadzenia przewodów szerokim łukiem.

Istotna jest także wygoda serwisu. Bardzo krótkie przewody pozwalają ograniczyć zajętą przestrzeń, ale podczas wymiany urządzenia muszą zostać odłączone i ponownie podłączone. Rozwiązanie, które dobrze wygląda w modelu 3D, może zatem utrudniać pracę monterowi i wydłużać obsługę urządzenia.

Wysokość urządzeń również wpływa na układ szafy. Najwyższy moduł zamontowany na szynie może wyznaczać wymagany prześwit od kanałów kablowych, a pośrednio także liczbę szyn możliwych do rozmieszczenia w danej przestrzeni. Dlatego warto oceniać nie tylko obrys samej

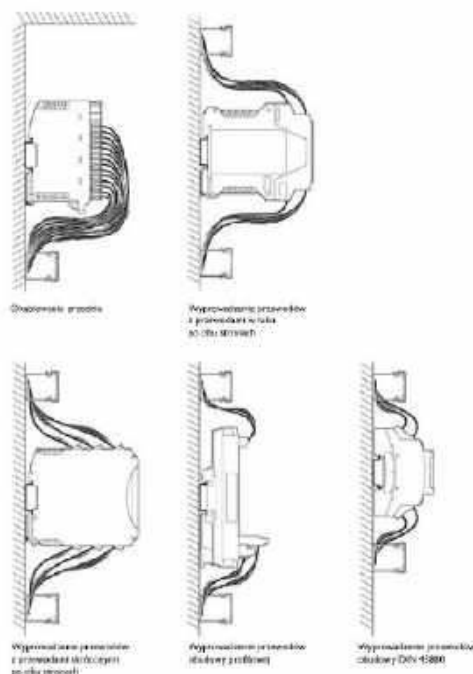


obudowy, lecz cały obszar montażowy potrzebny do jej okablowania, instalacji i późniejszej obsługi.

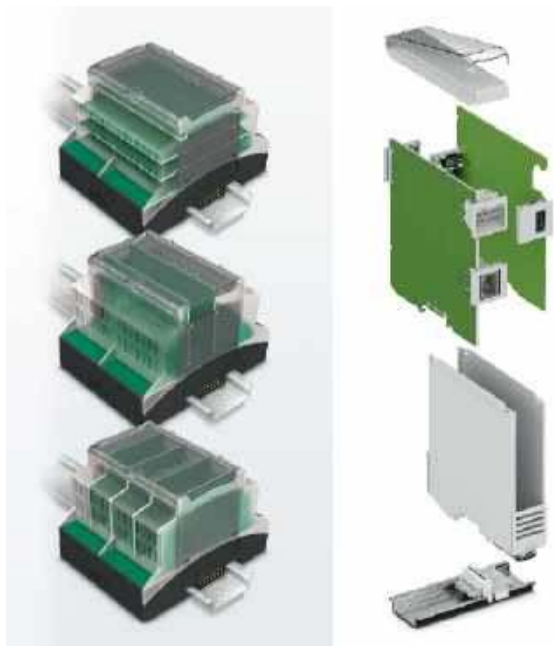
Konstrukcja obudowy wpływa na dostępną powierzchnię PCB

Wybór typu obudowy określa sposób montażu płytek PCB z elektroniką, ich orientację, dostępną powierzchnię montażową oraz położenie przyłączy. To z kolei ma wpływ na rozmieszczenie komponentów, wskaźników i interfejsów użytkownika.

W przypadku obudów modułowych prosta, dwuczęściowa budowa sprzyja szybkiemu montażowi końcowemu. Jeżeli jednak potrzebna jest duża liczba punktów zaciskowych umieszczonych na kilku poziomach, konstrukcja górnej



Prowadzenie przewodów vs. Konstrukcja obudowy



Możliwość montażu wielu płytek PCB

części obudowy może ograniczać miejsce dostępne na oznaczenia i elementy sygnalizacyjne.

Obudowy półskorupowe oferują inne możliwości. Pozwalają łatwiej zmieniać szerokość urządzenia i rozmieszczać przyłącza na różnych poziomach, bez tak dużego ograniczania powierzchni czołowej. Jednocześnie ich montaż końcowy jest bardziej złożony. Przy porównywalnych wymiarach zewnętrznych mogą natomiast zapewnić większy obszar montażowy PCB niż konstrukcje modułowe. Wybór staje się więc kompromisem pomiędzy dostępną przestrzenią dla elektroniki a prostotą składania urządzenia.

Jeszcze inną grupę stanowią konstrukcje profilowe, które pełnią przede wszystkim funkcję nośnika PCB. Profile mogą być docinane do wymaganej długości, a konstrukcja pozwala umieszczać płytki na kilku poziomach.

Nie istnieje zatem jeden rodzaj obudowy odpowiedni dla każdej elektroniki. Decyzja powinna wynikać z architektury urządzenia, liczby PCB, sposobu produkcji, wymaganej liczby przyłączy oraz potencjalnych czynności serwisowych.

Orientacja PCB i przyłącza powinny być projektowane razem

Poziomy lub pionowy montaż płytki względem szyny DIN nie jest wyłącznie kwestią mechaniczną. Orientacja PCB ogranicza wybór technologii przyłączeniowej i wpływa na kierunek prowadzenia przewodów. Poziomy układ płytki może zapewnić dużą powierzchnię czołową na wyświetlacze, wskaźniki i elementy sterujące. Nie oznacza jednak, że przewody muszą być prowadzone poziomo. Odpowiedni dobór przyłączy pozwala zmienić kierunek ich wyprowadzania, na przykład o 45° lub 90°, i dopasować go do warunków montażowych. W układzie pionowym często potrzebne są specjalnie dostosowane bloki zaciskowe. Dedykowane rozwiązania pozwalają zmienić

kierunek prowadzenia połączeń względem płytki i zapewnić wygodny dostęp.

Równie ważna jest decyzja pomiędzy przyłączem stałym a rozłącznym. Przyłącze stałe łączy przewód bezpośrednio z płytką i może oznaczać mniejszą liczbę elementów. Podczas serwisu przewody trzeba jednak odłączyć i podłączyć ponownie, co zwiększa znaczenie ich prawidłowego oznaczenia. W rozwiązaniu wtykowym okablowanie pozostaje połączone z wtyczką, dzięki czemu wymiana modułu może być prostsza, a zastosowanie kodowania ogranicza ryzyko błędnego ponownego połączenia.

Technologia przyłączeniowa, w kwestiach produkcyjnych, pozostaje nie bez mniejszego znaczenia. Rodzaj wybranego przyłącza wpływa na rodzaj montażu, lutowania na fali lub rozpliwowego, czy możliwość montażu automatycznego. Wtykowa konstrukcja może pozwolić na późniejszy wybór sposobu podłączania przewodów, ponieważ część przewodowa jest dodawana na dalszym etapie. Może to też stworzyć możliwość wyboru różnych rodzajów wtyków przez użytkowników końcowych.

To pokazuje, dlaczego obudowy, PCB i przyłączy nie należy rozwijać jako trzech oddzielnych obszarów projektu. Zmiana jednego z nich może wymusić korektę pozostałych.

Bilans cieplny

W zamkniętej obudowie komponenty elektryczne i mechaniczne oddają ciepło do otoczenia. Dopóki temperatury pozostają poniżej wartości granicznych, układ może działać prawidłowo. Po ich przekroczeniu konieczne jest zapewnienie skuteczniejszej wymiany ciepła, ponieważ nadmierna temperatura może prowadzić do uszkodzenia urządzenia.

Ciepło jest przekazywane przez promieniowanie, konwekcję i przewodnictwo. W przypadku obudów z tworzywa promieniowanie i przewodzenie przez ścianki mogą mieć ograniczoną skuteczność. Dodatkowo urządzenia zamontowane blisko siebie ograniczają oddawanie ciepła przez powierzchnie boczne, a sąsiednie moduły mogą stać się kolejnym źródłem nagrzewania. Zwiększenie odstępów poprawia warunki termiczne, ale jednocześnie zwiększa zapotrzebowanie na miejsce w szafie.



Zintegrowane interfejsy HMI



Dedykowane radiatory idealnie dopasowane do systemu obudów

Jeżeli obudowa ma szczeliny wentylacyjne i jest zamontowana w odpowiedniej pozycji, może wykorzystywać konwekcję naturalną. Chłodniejsze powietrze wpływa przez dolne otwory, ogrzewa się przy komponentach i opuszcza obudowę górą. Skuteczność tego mechanizmu zależy jednak od orientacji urządzenia i od tego, czy przepływ nie jest blokowany przez elementy znajdujące się wewnątrz lub w jego otoczeniu.

Na wynik cieplny wpływają zatem:

- moc tracona przez komponenty i ich rozmieszczenie na PCB,
- materiał i wielkość obudowy,
- pozycja montażowa,
- obecność otworów wentylacyjnych,
- sąsiednie urządzenia i kanały kablowe,
- sposób przyłączenia przewodów.

Symulacja termiczna umożliwia przygotowanie mapy cieplnej z uwzględnieniem elementów generujących ciepło oraz podzespołów wrażliwych na podwyższoną temperaturę. Pozwala sprawdzić wpływ ułożenia komponentów, obudowy i wentylacji jeszcze przed wykonaniem finalnej konstrukcji. W przypadkach zbliżonych do wartości granicznych wynik symulacji powinien zostać zweryfikowany pomiarem gotowego rozwiązania. Analiza może wskazać potrzebę użycia dodatkowych elementów chłodzących, pasywnych czy aktywnych, których wprowadzenie może okazać się problemowe na późniejszych etapach projektowania urządzenia.

Materiał obudowy trzeba dobierać do warunków pracy

Tworzywo obudowy nie pełni wyłącznie funkcji estetycznej. Znaczenie mają jego stabilność wymiarowa, zachowanie w określonej temperaturze, właściwości izolacyjne, odporność mechaniczna oraz wymagania dotyczące palności. Poszczególne materiały termoplastyczne różnią się właściwościami. Poliamid może być stosowany przy wyższych temperaturach, ale pochłania wodę i ma inną stabilność

wymiarową niż poliwęglan. Poliwęglan wyróżnia się sztywnością, większą odpornością mechaniczną, stabilnością wymiarową i ograniczonym pochłanianiem wilgoci. ABS znajduje zastosowanie tam, gdzie wymagana jest dobra odporność na uderzenia i jakość powierzchni.

Dobór materiału wymaga więc odniesienia do rzeczywistej aplikacji. Inne warunki występują w urządzeniu pracującym wewnątrz szafy, inne w obudowie zewnętrznej, a jeszcze inne w konstrukcji z komponentami generującymi dużo ciepła. Trzeba również sprawdzić zależność parametrów materiału od grubości ścianki. Wyższa odporność temperaturowa nie zawsze idzie w parze z najlepszą stabilnością wymiarową.

Trzy poradniki, jedna ścieżka projektowa

Znaczenie wyboru właściwej obudowy jest duże już na wstępnych etapach projektów. Wybrany system określa ilość dostępnego miejsca na PCB, możliwe kierunki podłączania przewodów, sposób montażu na szynie DIN i wygodę serwisowania. Dlatego Phoenix Contact przygotował trzy poradniki, które prowadzą przez kolejne obszary tego zagadnienia:

1. Obudowa elektroniki jako część szafy sterowniczej
Pokazuje zależności pomiędzy obudową, szyną DIN, technologią przyłączeniową, prowadzeniem przewodów i wykorzystaniem przestrzeni w szafie.
2. Konstrukcje obudów do elektroniki
Omawia konstrukcje modułowe, półskorupowe i profilowe, orientację PCB, przyłącza oraz dodatkowe elementy funkcjonalne urządzenia.
3. Rozpraszanie ciepła, materiały, testy
Porządkuje zagadnienia termiczne, kryteria wyboru tworzyw oraz metody weryfikacji właściwości materiału i kompletnej obudowy.

Materiały tworzą spójną bazę wiedzy dla konstruktorów projektujących urządzenia przeznaczone do montażu w szafach sterowniczych i innych instalacjach. Zamiast analizować pojedynczy parametr obudowy, pozwalają zobaczyć zależności pomiędzy mechaniką, elektroniką, produkcją i późniejszą eksploatacją. Zapisz się na serię trzech praktycznych poradników i poznaj najważniejsze kryteria wyboru obudów elektronicznych, od wymagań aplikacji po personalizację gotowego rozwiązania. Znajdziesz je na <https://www.phoenixcontact.com/pl-pl/zawsze-wlasciwa-obudowa/poradnik-obudowy-formularz>.

Więcej informacji: Phoenix Contact
Paweł Zientarski
Menedżer Obszaru Biznesu
– złącza i obudowy do elektroniki
tel. +48 694 485 087
e-mail: pzientarski@phoenixcontact.pl



Obudowy w elektronice profesjonalnej

W konstrukcji amatorskiej obudowa bywa ostatnim etapem pracy: układ działa, więc trzeba go w coś zamknąć. W elektronice profesjonalnej kolejność jest odwrotna. Obudowa jest komponentem, który wybiera się na początku, bo narzuca wymiary płytki, sposób prowadzenia przewodów, dostępny budżet cieplny, osiągalny stopień ochrony i – co bywa najbardziej kosztowne, jeśli zostanie przeoczone – ścieżkę certyfikacji wyrobu. Ten artykuł jest przeglądem asortymentu i kryteriów doboru obudów dostępnych konstruktorowi działającemu w Polsce.

Obudowa jest komponentem, nie opakowaniem

Obudowa urządzenia profesjonalnego pełni równocześnie kilka funkcji technicznych i każda z nich generuje osobny zestaw wymagań. Po pierwsze jest elementem konstrukcji mechanicznej: przenosi obciążenia z gniazd i przewodów na punkty mocowania, ustala położenie płytki względem panelu czołowego, chroni układ przed drganiami i uderzeniami. Po drugie jest barierą środowiskową – decyduje, ile pyłu, wilgoci i promieniowania nadfioletowego dotrze do elektroniki w ciągu kilkunastu lat eksploatacji. Po trzecie bierze udział w bilansie cieplnym, bo to przez jej ścianki ciepło strat opuszcza urządzenie. Po czwarte, w urządzeniach zasilanych z sieci, jest elementem ochrony przeciwporażeniowej i przeciwpożarowej, którego właściwości podlegają ocenie w procesie certyfikacji. Po piąte wreszcie stanowi interfejs użytkownika i nośnik wzornictwa.

Rzadko się zdarza, aby te pięć funkcji dało się optymalizować niezależnie. Otwory wentylacyjne poprawiające bilans cieplny obniżają stopień ochrony. Metalowa ścianka,

dobra jako ekran elektromagnetyczny, wymaga izolowania płytki i przemyslenia uziemienia. Uszczelka o dużej sile docisku poprawia szczelność, ale wymusza grubsze ścianki i sztywniejszą pokrywę. Dobór obudowy jest więc w praktyce serią kompromisów, a ich jakość zależy od tego, jak wcześniej konstruktor zacznie je świadomie rozstrzygać.

Co pokazują dane o rynku krajowym

Zanim przejdziemy do techniki, warto spojrzeć na to, czego konstruktor może się na polskim rynku spodziewać. Corocznie prowadzone badanie ankietowe wśród producentów i dystrybutorów obudów działających w Polsce, publikowane w Informatorze Rynkowym Elektroniki, daje w tej sprawie kilka konkretnych wskazówek.

Na pytanie o to, na jakie typy obudów jest największy zbyt na rynku krajowym, 68% ankietowanych dostawców wskazało małe obudowy z tworzyw sztucznych. Drugą pozycję, z wynikiem 48%, zajęły obudowy na szynę DIN, a trzecią – ex aequo z dużymi szafami przemysłowymi – małe obudowy metalowe, wskazane przez 45% dostawców. Duże obudowy plastikowe wskazało zaledwie 13% ankietowanych. Wniosek dla konstruktora jest praktyczny: w segmencie małych obudów z tworzyw i obudów na szynę DIN oferta katalogowa jest najgłębsza, dostępność najlepsza, a szanse znalezienia wyrobu pasującego bez modyfikacji – największe.

Drugie pytanie ankiety dotyczyło cech oferty, którymi klienci kierują się przy wyborze dostawcy obudów. Najczęściej wskazywanym kryterium – przez 77% ankietowanych – był termin wykonania lub dostawy. Niska cena znalazła się tuż za nim, z wynikiem 74%. Wysoką jakość i dokładność wykonania wskazało 55% badanych, dostępność produktów na zamówienie 52%, a usługi indywidualizacji wersji seryjnych 45%. To układ charakterystyczny dla rynku, na którym terminowość stała się de facto kryterium czasu wyprzedziło kryterium ceny,



Fotografia 1. Obudowa obiektowa zamontowana na konstrukcji zewnętrznej – przykład aplikacji, w której funkcja bariery środowiskowej dominuje nad pozostałymi

Na które typy obudów jest największy zbyt na rynku krajowym?



Rysunek 1. Struktura zbytu obudów na rynku krajowym według badania IRE 2026 – udział wskazań ankietowanych dostawców dla poszczególnych kategorii wyrobów

a dostępność wersji modyfikowanej znalazła się wyżej niż kiedykolwiek wcześniej.

Trzecie zestawienie dotyczyło trendów zmieniających rynek obudów. Usługi indywidualizacji i obróbki mechanicznej wskazało 68% ankietowanych dostawców, wysoką dokładność wykonania i jakość 48%, uniwersalność i ergonomię produktów również 48%, a modułowość i skalowalność konstrukcji 42%. Możliwość instalacji wielu płytek i złączy oraz prostotę mocowania komponentów wskazało po 39% badanych, natomiast dobre zarządzanie ciepłem – 26%.

Ostatnia liczba zasługuje na komentarz. To, że gospodarka ciepła obudowy znalazła się nisko w rankingu trendów rynkowych, nie oznacza, że jest problemem marginalnym. Oznacza raczej, że dostawcy rzadko słyszą o niej od klientów na etapie zakupu – a to zwykle znaczy, że temat może wracać później, w postaci przegrzewającego się prototypu.

Mapa asortymentu

Podział obudów jest trudny, bo kryteria się przenikają: ta sama obudowa może być jednocześnie plastikowa, modułowa i przeznaczona na szynę. Dla celów doboru najwygodniejszy jest podział według sposobu montażu, bo to on w pierwszej kolejności ogranicza listę kandydatów.

Obudowy na szynę DIN

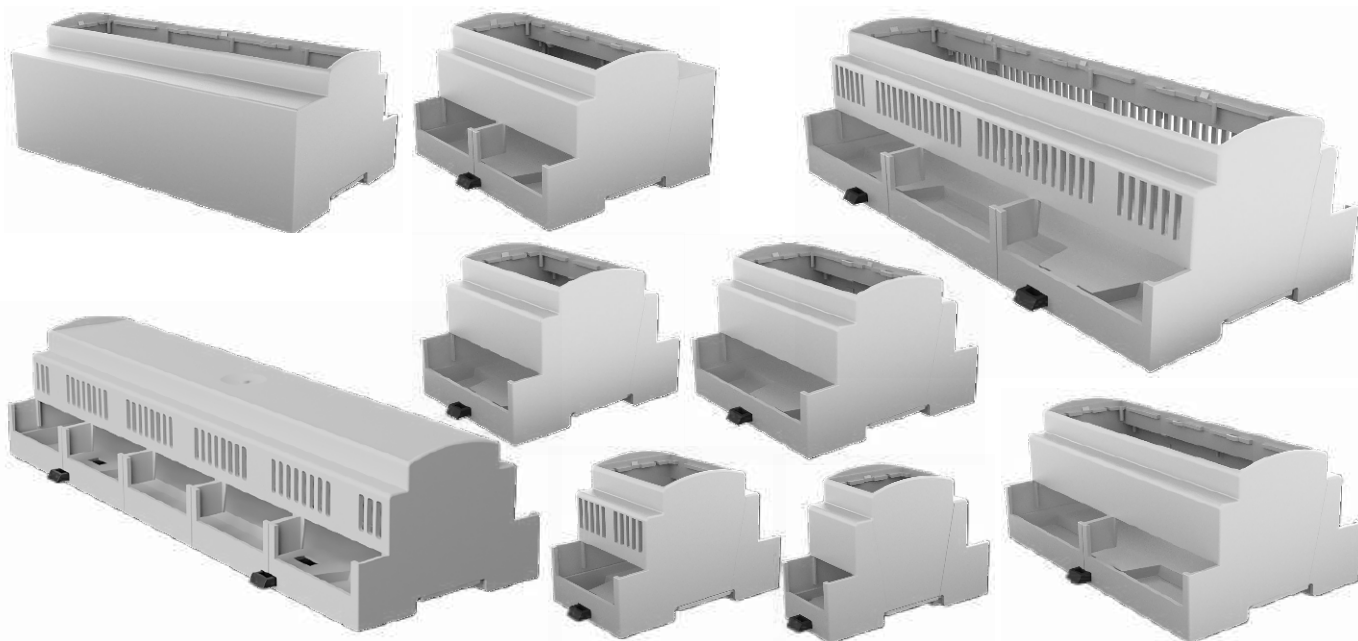
To najbardziej ustandaryzowana rodzina w całym asortymencie i zarazem ta, w której konstruktor ma największą swobodę. Szyna montażowa TS35 opisana normą EN 60715 jest obecna w każdej szafie sterowniczej i w każdej rozdzielnicy instalacyjnej, a obudowy do niej pasujące wykonuje kilkudziesięciu producentów.

Obudowy na szynę dzielą się na dwie podgrupy o różnym rodowodzie. Pierwsza to wyroby zgodne z niemiecką

normą DIN 43880, opisującą wymiary gabarytowe aparatury instalacyjnej. Norma ta definiuje moduł szerokości równy 17,5 mm oraz wysokości pozwalające zmieścić urządzenie za listwą maskującą rozdzielnicę. Obudowa dwumodułowa ma zatem 35 mm szerokości, czteromodułowa 70 mm i tak dalej. Zgodność z DIN 43880 oznacza, że urządzenie zmieści się w typowej rozdzielnicy mieszkaniowej lub obiektowej obok wyłączników nadprądowych, nie wystając ponad ich lico. Dla urządzeń automatyki budynkowej ta zgodność jest praktycznie warunkiem wejścia na rynek.

Druga podgrupa to obudowy przemysłowe, przeznaczone do szaf sterowniczych, w których nie ma listwy maskującej i wymiar wysokości nie jest tak sztywno narzucony. Obudowy te bywają wyższe i głębsze, mają najczęściej kształt zbliżony do kielicha – dolna część z nóżką zatrzaskową na szynę, część środkowa mieszcząca płytkę i część górna z zaciskami. Typowe szerokości to 17,6 mm, 22,5 mm, 35 mm i 45 mm; wysokości rzędu 99...100 mm, a głębokości 105...115 mm. Wersje ze szczelinami wentylacyjnymi i bez nich występują zwykle równolegle w tej samej serii.

Cechą, która odróżnia współczesne obudowy na szynę od wyrobów sprzed dekady, jest złącze magistralowe montowane na szynie pod obudową. Jest to listwa z kilkoma połączonymi stykami, wpinana w szynę przed zatrzasknięciem obudów; sąsiadujące moduły łączą się przez nią automatycznie, bez ani jednego przewodu w szafie. Spotykane rozwiązania mają od 5 do 16 styków i pozwalają rozprzecznić między modułami zasilanie oraz sygnał magistrali szeregową, przy obciążalności rzędu kilku amperów i napięciu do 125 V. Część konstrukcji przewiduje styki równoległe, dzięki którym wyjęcie jednego modułu nie przerywa połączenia w pozostałych – to istotne, gdy urządzenia mają być wymieniane bez wyłączania instalacji.



Fotografia 2. Obudowy wielomodułowe na szynie DIN wykonane zgodnie z DIN 43880, z panelami czołowymi i pokrywami zaciskowymi

Dla konstruktora złącze magistralowe na szynie oznacza konkretną decyzję projektową: podział urządzenia na moduły staje się opłacalny. Zamiast jednej dużej płytki w obudowie ośmiomodułowej można wykonać rodzinę modułów dwumodułowych, skonfigurowaną u klienta.

Obudowy modułowe i systemy obudów

Modułowość to pojęcie, którego znaczenie w ofertach producentów uległo zawężeniu. Kiedyś oznaczała spójność całego katalogu jednego producenta; dziś dotyczy zwykle jednej serii albo rodziny, w obrębie której poszczególne detale są wymienne.

System obudów w rozumieniu współczesnych katalogów to zestaw części składowych – podstaw, elementów pośrednich, pokryw, paneli czołowych, uchwytów płytek, przepustów, nóżek, zaślepek – z których buduje się obudowę o wymaganej wysokości, szerokości i wyposażeniu. Analogia do klocków jest tu naprawdę na miejscu. Zysk jest podwójny: konstruktor dostaje

wymiar dopasowany do płytki bez zamawiania narzędzi, a producent urządzenia zachowuje spójną koncepcję wzorniczą dla całej rodziny wyrobów.

Ograniczeniem systemów modułowych jest siatka wymiarów. Elementy składowe występują w skończonej liczbie rozmiarów i jeśli płytka jest o 3 mm za szeroka, trzeba albo przeprojektować płytkę, albo przejść o jeden rozmiar wyżej i pogodzić się z pustą przestrzenią. Dlatego wybór systemu obudów powinien nastąpić przed wytrasowaniem obwodu drukowanego, a nie po nim.

Obudowy dla komputerów jednopłytkowych

Osobno warto odnotować kategorię, która powstała w ostatnich latach i jest dla czytelników EP szczególnie interesująca: obudowy katalogowe przeznaczone dla popularnych komputerów jednopłytkowych. Producenci systemów obudów wprowadzili do oferty warianty z gotowym rozstawem otworów montażowych i wycięciami pod złącza tych płytek, w wykonaniach na szynie DIN zgodnych z DIN 43880 oraz w wykonaniach obiektowych do montażu poza szafą, uzupełnionych o podstawki i adaptory biurkowe. Część z nich przewiduje złącze magistralowe na szynie, przez które komputer jednopłytkowy komunikuje się z modułami wejść i wyjść.

Znaczenie tej kategorii wykracza poza wygodę. Komputer jednopłytkowy w obudowie przemysłowej, z deklarowanym stopniem ochrony i z uporządkowanym przyłączem, przestaje być rozwiązaniem prototypowym, a staje się elementem, który da się zainstalować u klienta i objąć serwisem. Dla małych producentów urządzeń jest to często najkrótsza droga od koncepcji do wyrobu.



Fotografia 3. Modułowe obudowy na szynie DIN z metalową nóżką zatraskową i złączem magistralowym montowanym na szynie pod obudowami



Fotografia 4. Obudowy na szynę DIN i obiektowe przeznaczone dla komputerów jednopłytkowych, z wycięciami pod złącza płytki



Fotografia 5. Obudowa hermetyczna z poliwęglanu z przezroczystą pokrywą i uszczelką obwodową, w wykonaniu IP65/IK07



Fotografia 6. Obudowy ręczne i pulpitowe z zagłębieniem pod klawiaturę membranową i komorą baterii

Obudowy naścienne i obiektowe

Kategoria obudów obiektowych – w katalogach spotyka się też określenia „polowe” i „terenowe” – obejmuje wyroby przeznaczone do montażu poza szafą: na ścianie, na maszcie, na konstrukcji wsporczej, na barierce, w otwartej przestrzeni hali produkcyjnej lub na zewnątrz budynku.

Wymagania są tu zupełnie inne niż w szafie. Obudowa obiektowa musi znosić bezpośrednie nasłonecznienie, opady, mgłę solną w aplikacjach nadmorskich, cykliczne zmiany temperatury i wynikającą z nich kondensację. Typowe wykonania to poliwęglan wzmacniany, odporny na promieniowanie nadfioletowe, o grubości ścianek dochodzącej do 4 mm, z uszczelką wylewaną lub zintegrowaną, w wykonaniach o stopniu ochrony od IP66 wzwyż. Zakresy temperatur pracy deklarowane dla lepszych serii sięgają od -40°C do $+80^{\circ}\text{C}$, z krótkotrwałą odpornością na temperatury wyższe.

Cechy, na które warto zwrócić uwagę przy porównywaniu ofert w tej kategorii, to: dostępność wersji z pokrywą przezroczystą (przydatna, gdy w środku jest wyświetlacz lub sygnalizacja), możliwość plombowania, zintegrowane zabezpieczenie przed nieuprawnionym otwarciem, zintegrowane regulowane mocowanie płytki lub panelu oraz – rzecz często pomijana – membrana wyrównująca ciśnienie.

Membrana wyrównująca ciśnienie zasługuje na osobne zdanie, bo jej brak jest jedną z częstszych przyczyn awarii urządzeń zewnętrznych. Szczelna obudowa nagrzewająca się w słońcu i chłodzona nocą działa jak pompa: wewnątrz na przemian rozpręża się i kurczy, a przy każdym cyklu przez najsłabsze miejsce uszczelnienia zasysane jest wilgotne powietrze. Po kilkudziesięciu cyklach w obudowie zbiera się woda, choć formalnie ma ona stopień ochrony IP67. Membrana z materiału przepuszczającego parę wodną, ale nie wodę ciekłą, wyrównuje ciśnienie i przerywa ten mechanizm. W obudowach przeznaczonych na zewnątrz jest wyposażeniem, o które warto zapytać wprost.

Obudowy tablicowe, pulpitemowe i ręczne

Obudowy tablicowe montuje się w wycięciu płyty czołowej szafy lub pulpitu operatorskiego. Wymiary wycięć są w tej kategorii częściowo ustandaryzowane, co ułatwia wymianę urządzeń różnych producentów. Uszczelnienie realizowane jest zwykle od strony czołowej – stąd typowa deklaracja producenta, że urządzenie ma IP65 od frontu i IP20 od tyłu.

Obudowy pulpitemowe, o pochylonym panelu czołowym, przeznaczone są do urządzeń stołowych obsługiwanych przez operatora. Obudowy ręczne, w tym te z komorą baterii i zagłębieniem pod klawiaturę foliową, tworzą osobną, dużą grupę wyrobów, w której obok kryteriów technicznych



Obudowy nigdy dotąd nie były tak uniwersalne

Modułowe obudowy ICS dla nowoczesnych urządzeń IoT.

- Wysoka elastyczność zastosowania dzięki systemowi modułowemu
- Zintegrowana technika przyłączeniowa
- Optymalne wykorzystanie przestrzeni oraz możliwość dostosowania konstrukcji, kolorów i nadruku
- 8-biegunowe łączniki T-BUS na szynę DIN do łatwej komunikacji między modułami



Odwiedź naszą stronę internetową i dowiedz się więcej!

www.phoenixcontact.pl/Obudowy





Fotografia 7. Obudowa na szynę DIN z aluminium radiatorem zintegrowanym z korpusem, przeznaczona do modułów o podwyższonej mocy strat

naprawdę liczy się ergonomia: kształt dopasowany do dłoni, rozkład masy, umiejscowienie punktów chwytu.

Obudowy z profili aluminiowych i obudowy z radiatorem

Konstrukcja oparta na profilu wyciskany łączy trzy cechy jednocześnie: dużą sztywność, dobre ekranowanie i zdolność odprowadzania ciepła. Profil stanowi ścianki boczne i sufit, a płyta czołowa i tylna zamykają obudowę

od przodu i tyłu. Płytkę wsuwa się w rowki wykonane w profilu, co eliminuje kołki dystansowe i ułatwia montaż.

Odmianą tej konstrukcji, ważną dla układów mocy, są obudowy na szynę DIN z radiatorem zintegrowanym z korpusem. Spotyka się dwie topologie: sam profil wy-ciskany pełniący funkcję radiatora oraz układ z rozpra-szczaczem ciepła (heat spreader) przekazującym ciepło z elementów mocy na profil. Deklarowane szerokości ra-diatora w wyrobach katalogowych mieszczą się w zakresie



Fotografia 8. Kaseta 19-calowa z wentylatorami i wysuwaną płytą montażową – konstrukcja typowa dla aparatury instalowanej w szafie rack

od ok. 9 mm do ok. 54 mm, przy maksymalnej wysokości elementów montowanych po stronie radiatora rzędu 10...12 mm. Część producentów udostępnia dla tych serii wyniki symulacji termicznej, co pozwala oszacować dopuszczalną moc strat jeszcze przed budową prototypu.

Obudowy 19-calowe, szafy i wykonania specjalne

Kaseta 19-calowa i szafa rack pozostają standardem tam, gdzie urządzenie ma współistnieć z innym sprzętem: w telekomunikacji, aparaturze pomiarowej, systemach teleinformatycznych. Podział na jednostki wysokości (1U=44,45 mm) i szerokości modułu (1HP=5,08 mm) jest tu na tyle utrwalony, że w praktyce projektowej traktuje się go jak dane wejściowe.

Duże szafy przemysłowe stanowią odrębną kategorię produktową, konstrukcyjnie różniącą się od pozostałych. Buduje się je na szkielecie z profili metalowych, do którego dołącza się ściany, podłogę i sufit w wykonaniach dobieranych przez użytkownika. Materiałem jest zwykle blacha stalowa, często fosforanowana i malowana proszkowo. Szafy rzadko mają z góry zdefiniowaną konstrukcję – użytkownik komponuje je z elementów składowych, dobierając wielkość, sposób łączenia, materiał drzwi oraz wyposażenie dodatkowe: wentylatory, klimatyzatory, panele oświetleniowe, systemy rozdziału energii.

Na obrzeżach asortymentu znajdują się wykonania specjalne: obudowy w wykonaniu przeciwwybuchowym (ze stali nierdzewnej lub z poliestru, z odpowiednim dopuszczeniem dla stref zagrożonych wybuchem), obudowy higieniczne dla przemysłu spożywczego,

obudowy o podwyższonej odporności dla taboru kolejowego. Są to produkty niszowe, ale dostępne u dostawców obecnych na rynku krajowym, a ich wspólną cechą jest to, że nie kupuje się ich na podstawie samego rysunku wymiarowego – potrzebna jest analiza dopuszczeń.

Materiały i to, co z nich wynika

Podstawowy podział materiałowy – metale i tworzywa sztuczne – jest prosty, ale konsekwencje wyboru sięgają daleko.

Wśród tworzyw najczęściej stosowane są ABS, polistyren, poliamid, poliwęglan oraz mieszanki i wersje wzmacniane, zwykle włóknem szklanym. Różnią się przede wszystkim odpornością na temperaturę i narażenia mechaniczne. Poliwęglan wypada dobrze w obu tych kategoriach i dodatkowo, w odmianach stabilizowanych, dobrze znosi promieniowanie nadfioletowe, dlatego zdominował segment obudów zewnętrznych. ABS jest tańszy i łatwiejszy w obróbce wykończeniowej, ale gorzej znosi ekspozycję słoneczną – w zastosowaniach zewnętrznych bez dodatkowej ochrony żółknie i kruszeje.

Dla urządzeń podlegających certyfikacji istotna jest klasyfikacja palności według UL 94. W praktyce w aparaturze przemysłowej i instalacyjnej oczekuje się klasy V-0, a coraz częściej – dla obudów pełniących funkcję obudowy ognioochronnej – także określonych właściwości w badaniu rozżarzonym drutem, opisywanych wskaźnikami GWFI i GWIT. Warto sprawdzić te dane w karcie katalogowej obudowy, zanim urządzenie trafi do laboratorium.

Metale to przede wszystkim blacha stalowa i aluminium. Stal zapewnia dużą wytrzymałość mechaniczną i jest materiałem najtańszym w kategorii obudów większych, dlatego dominuje w skrzynkach instalacyjnych i szafach. Aluminium – w postaci blachy giętej, odlewu ciśnieniowego lub profilu wyciskanego – jest lżejsze i lepiej przewodzi ciepło, dlatego przeważa w obudowach urządzeń, w których liczy się bilans cieplny lub masa. Wykończenie powierzchni ma tu znaczenie nie tylko estetyczne: malowanie proszkowe daje trwałą powłokę odporną na czynniki środowiskowe, natomiast anodowanie aluminium wytwarza warstwę tlenku, która jest izolatorem – co bywa



Fotografia 9. Skrzynka przyłączeniowa ze stali nierdzewnej w wykonaniu przeciwwybuchowym, z dopuszczeniem ATEX i IECEx

Rozżarzony drut – po co ten test

Badanie rozżarzonym drutem (glow wire) polega na przyłożeniu do próbki materiału rozgrzanego drutu o zadanej temperaturze i obserwacji, czy próbka się zapali oraz jak szybko gaśnie po odsunięciu drutu. GWFI (Glow Wire Flammability Index) określa najwyższą temperaturę, przy której materiał danej grubości nie pali się dłużej niż 30 s po odsunięciu drutu. GWIT (Glow Wire Ignition Temperature) to temperatura o 25 K wyższa od najwyższej, przy której materiał się nie zapalił. Normy wyrobu odwołują się do tych wskaźników, gdy obudowa ma powstrzymać rozprzestrzenienie ognia zainicjowanego wewnątrz urządzenia – na przykład przez przegrzany element.

zaskoczeniem dla konstruktora liczącego na ciągłość galvaniczną między częściami obudowy.

Warto odnotować kierunek, w jakim rynek się przesuwają. Miniaturyzacja i rosnący udział urządzeń przenośnych sprzyjają tworzywom. Obudowy metalowe pozostają w niszach wymagających największej trwałości, ekranowania elektromagnetycznego lub odporności na promieniowanie nadfioletowe: w urządzeniach telekomunikacyjnych pracujących na zewnątrz, w taborze kolejowym i morskim, w rozległych instalacjach przemysłowych, w urządzeniach radiowych. Równocześnie hermetyczne obudowy z tworzyw o poprawionej wytrzymałości mechanicznej, z uszczelkami wlewianymi i z poliwęglanu wzmacnianego włóknem, zbliżyły się funkcjonalnie do wersji metalowych na tyle, że w wielu aplikacjach są rozwiązaniem równorzędnym.

Interfejs obudowa-płytki

To miejsce, w którym najczęściej rodzą się problemy, bo styka się tu praca dwóch osób lub dwóch działów: projektanta obwodu drukowanego i konstruktora mechanika.

Punktem wyjścia jest sposób ustalenia położenia płytki. Spotyka się cztery rozwiązania. Mocowanie na kołkach dystansowych, z otworami w płytce – najbardziej uniwersalne, ale wymagające precyzyjnego zwymiarowania otworów już na etapie projektu obwodu. Prowadnice w ściankach, w które płytkę się wsuwa – szybkie w montażu, wymagające dokładnego zwymiarowania krawędzi płytki i zwykle uniemożliwiające umieszczenie elementów przy samej krawędzi. Uchwyty płytki będące osobnym detalem systemu obudów, przesuwane w profilu i pozwalające zamocować płytkę o dowolnej długości, a nierzadko także kilka płytek jedna nad drugą. Wreszcie mocowanie na wypustach (domach) w dolnej części obudowy, wykorzystywane w wyrobach tanich.



Fotografia 10. Płyta montażowa i tuleje dystansowe jako sposób ustalenia położenia płytki wewnątrz obudowy

Drugie zagadnienie to złącza. W obudowach na szynę DIN część górna obudowy jest zwykle projektowana razem z zaciskami: producent oferuje gotowe pokrywy z wbudowanymi terminalami do PCB, dobranymi do rastra i wysokości montażu. Warto zwrócić uwagę na technikę przyłączeniową zastosowaną w tych zaciskach, bo przekłada się ona bezpośrednio na czas montażu u klienta końcowego. Technika Push-in, w której przewód zakończony tulejką wsuwa się bez użycia narzędzia, jest dziś standardem w aparaturze instalacyjnej i wypiera połączenia śrubowe.

Trzecie zagadnienie to tolerancje. Obudowa wtryskiwana ma tolerancje wymiarów rzędu dziesiątych części milimetra i podlega skurczowi po wyjęciu z formy; płytka obwodu drukowanego wykonana metodą frezowania konturu ma tolerancje mniejsze, ale niezerowe. Jeśli konstrukcja zakłada, że gniazdo zamontowane na płytce ma dokładnie trafić w otwór w panelu czołowym, kumulacja tolerancji szybko staje się problemem. Rozwiązania są dwa: albo daje się gniazdu luz i maskuje go elementem panelu, albo pozycję gniazda ustala panel, a płytka ma swobodę. To drugie podejście, powszechne w konstrukcjach profesjonalnych, wymaga jednak, aby połączenie płytki z gniazdem było elastyczne.

Ciepło

Bilans cieplny urządzenia w obudowie zamkniętej sprowadza się do prostego stwierdzenia: cała moc strat musi opuścić obudowę przez jej powierzchnię zewnętrzną, a w warunkach ustalonych przyrost temperatury wnętrza ponad otoczenie jest w przybliżeniu proporcjonalny do tej mocy.

Praktyczna konsekwencja jest taka, że o dopuszczalnej mocy strat decyduje przede wszystkim pole powierzchni obudowy i to, jak jest ona zorientowana. Obudowa stojąca pionowo, z drogą swobodnej konwekcji wzdłuż ścianek,



Fotografia 11. Pokrywy zaciskowe i panele czołowe jako wymienne elementy systemu obudów modułowych

odprowadza wyraźnie więcej ciepła niż ta sama obudowa leżąca płasko na stole. Obudowa wciśnięta w szafę między dwie inne, o temperaturze otoczenia podwyższonej do 50°C, ma odpowiednio mniejszy zapas.

Konstruktor ma do dyspozycji kilka narzędzi. Otwory wentylacyjne – najtańsze i najskuteczniejsze, ale ograniczające stopień ochrony do IP20 lub IP30 i wprowadzające drogę dla pyłu. Radiator zintegrowany z obudową, opisany wcześniej, przenoszący ciepło z elementów mocy bezpośrednio na powierzchnię zewnętrzną. Przekładki termoprzewodzące łączące elementy mocy z metalową ścianką obudowy – rozwiązanie wygodne, ale wymagające uwzględnienia w projekcie tolerancji i sił docisku, a przy ściankach uziemionych także wymagań izolacyjnych. Wreszcie wentylator, którego zastosowanie zmienia klasyfikację urządzenia pod względem niezawodności i wymaga przewidzenia filtra oraz jego wymiany.

W przypadku obudów katalogowych warto pytać dostawcę o dostępność wyników symulacji termicznej. Część producentów systemów obudów udostępnia takie dane dla swoich serii radiatorowych, podając zależność przyrostu temperatury od mocy strat dla różnych szerokości profilu. Dysponując tymi danymi, można wyeliminować oczywiste błędy doboru bez budowy modelu.

Szczelność, odporność mechaniczna i kondensacja

Stopień ochrony obudowy opisuje kod IP zdefiniowany w normie PN-EN 60529. Kod ten jest w branży powszechnie znany, ale bywa interpretowany zbyt swobodnie, dlatego przypominamy jego strukturę.

Odporność na uderzenia opisuje osobny kod IK, zdefiniowany w normie PN-EN 62262. Jest to dwucyfrowe oznaczenie od IK00 do IK10, odpowiadające energii uderzenia

Tabela 1. Stopnie ochrony IP według PN-EN 60529 – znaczenie pierwszej i drugiej cyfry charakterystycznej

Pierwsza cyfra charakterystyczna – ochrona przed ciałami stałymi i dostępem do części niebezpiecznych		
Cyfra	Ochrona przed ciałami stałymi	Ochrona przed dostępem
0	brak	brak
1	ciała o średnicy ≥ 50 mm	wierzchem dłoni
2	ciała o średnicy $\geq 12,5$ mm	palcem
3	ciała o średnicy $\geq 2,5$ mm	narzędziem
4	ciała o średnicy ≥ 1 mm	drutem
5	pył w stopniu nie pogarszającym działania	drutem
6	pyłoszczelność	drutem
Druga cyfra charakterystyczna – ochrona przed wodą		
Cyfra	Rodzaj narażenia	
0	brak ochrony	
1	padające krople wody	
2	krople przy wychyleniu obudowy do 15° od pionu	
3	natrysk pod kątem do 60° od pionu	
4	bryzgi z dowolnego kierunku	
5	struga wody 12,5 l/min z dowolnej strony	
6	silna struga wody 100 l/min z dowolnej strony	
7	krótkotrwałe zanurzenie (30 min, 0,15 m ponad wierzch obudowy)	
8	zanurzenie ciągłe w warunkach uzgodnionych z producentem	
9	strumień o wysokim ciśnieniu i podwyższonej temperaturze	

Jak czytać kod IP

Pierwsza cyfra charakterystyczna (0...6) opisuje ochronę przed dostępem do części niebezpiecznych i przed wnikaniami ciał stałych: 4 to ochrona przed ciałami o średnicy 1 mm i przed dostępem drutem, 5 to ochrona przed pyłem w stopniu nie pogarszającym działania urządzenia, 6 to pełna pyłoszczelność. Druga cyfra (0...9) opisuje ochronę przed wodą: 4 to bryzgi z dowolnego kierunku, 5 to struga wody o wydatku 12,5 l/min, 6 to silna struga o wydatku 100 l/min, 7 to krótkotrwałe zanurzenie (30 min na głębokość 0,15 m ponad wierzch obudowy), 8 to zanurzenie ciągłe w warunkach uzgodnionych między producentem a użytkownikiem, 9 to strumień wody o wysokim ciśnieniu i podwyższonej temperaturze. Dwie rzeczy warto zapamiętać. Po pierwsze, wyższa druga cyfra nie zawsze zawiera w sobie niższą: IP67 nie oznacza automatycznie spełnienia wymagań IP65, bo badania są różne – dlatego producenci deklarują niekiedy dwa kody równolegle, w postaci IP66/67. Po drugie, kod IP opisuje wynik badania nowej obudowy w warunkach laboratoryjnych, a nie stan po pięciu latach ekspozycji uszczelki na promieniowanie nadfioletowe.



Fotografia 12. Obudowa ścienna ze stali nierdzewnej z uszczelnieniem z pianki poliuretanowej nakładanej na obwódzie

od poniżej 0,14 J do 20 J. W obudowach przeznaczonych do miejsc publicznych lub narażonych na uszkodzenia mechaniczne deklaracja IK bywa równie ważna jak IP, a bywa pomijana w porównaniach ofert.

Osobnego omówienia wymaga technika uszczelnienia, bo różnice między rozwiązaniami przekładają się na koszt i na trwałość. Uszczelka wkładana, najczęściej w postaci sznura lub oringu w rowku, jest tania i wymienna, ale wrażliwa na dokładność montażu. Uszczelka wylewana – w katalogach spotyka się skrót FIPFG, od Formed In Place Foam Gasket – nakładana jest robotem w postaci pianki poliuretanowej, która polimeryzuje bezpośrednio na krawędzi obudowy. Daje ciągłość uszczelnienia na całym obwodzie, także w narożnikach, i pozwala uzyskać wysokie stopnie ochrony w obudowach z tworzyw. Jej wadą jest to, że przy uszkodzeniu nie da się jej wymienić w warunkach serwisowych.

Wróćmy jeszcze do kondensacji, bo jest to problem, którego kod IP nie opisuje w ogóle. Wilgoć wewnątrz szczelnej obudowy pochodzi z dwóch źródeł: z powietrza zamkniętego w niej podczas montażu oraz z wnikania przez cykle ciśnieniowe. Środki zaradcze to membrana wyrównująca ciśnienie, montaż w kontrolowanej wilgotności,

a w wykonaniach krytycznych – pochłaniacz wilgoci lub zalanie układu masą. Zalanie obudowy ma przy okazji dwie inne zalety, o których warto wiedzieć: ogranicza możliwość nieuprawnionej ingerencji w układ i poprawia odprowadzanie ciepła z elementów do ścianek.

Obudowa jako element bezpieczeństwa wyrobu

W urządzeniach zasilanych z sieci obudowa przestaje być kwestią wyboru konstruktora, a staje się przedmiotem wymagań normatywnych. Dwie normy zbiorcze mają tu zastosowanie najczęściej: PN-EN IEC 62368-1 dla sprzętu elektronicznego audio, wideo, teleinformatycznego i telekomunikacyjnego oraz PN-EN 61010-1 dla elektrycznych przyrządów pomiarowych, automatyki i urządzeń laboratoryjnych.

Z punktu widzenia obudowy wymagania sprowadzają się do kilku grup. Obudowa ma uniemożliwiać dostęp do części czynnych – sprawdzany jest on znormalizowanym palcem probierczym i sondą, a w niektórych przypadkach także dostęp narzędziem. Ma zachować swoje właściwości ochronne po badaniach mechanicznych: próbie zrzućtu, nacisku, uderzenia. Ma pełnić funkcję obudowy ognioochronnej, czyli powstrzymać rozprzestrzenienie ognia zainicjowanego wewnątrz – stąd wymagania dotyczące klasy palności materiału oraz konstrukcji otworów wentylacyjnych, których geometria i rozmieszczenie względem elementów mogących się zapalić są w normach opisane wprost. Ma wreszcie zapewnić wymagane odstępstwa izolacyjne, powierzchniowe i w powietrzu, między obwodami o różnych napięciach oraz między obwodami a dostępnymi częściami metalowymi.

Praktyczna wskazówka jest taka: jeśli urządzenie ma być zasilane z sieci, warto wybrać obudowę, dla której producent deklaruje właściwości materiałowe potrzebne w certyfikacji i udostępnia odpowiednie dokumenty. Obudowa kupiona bez tych danych może się okazać najdroższym elementem projektu – nie ceną, lecz czasem straconym w laboratorium.

Ekranowanie i kompatybilność elektromagnetyczna

Obudowa metalowa może być ekranem, ale nie jest nim automatycznie. O skuteczności ekranowania decydują nie ścianki, lecz nieciągłości: szczeliny między częściami, otwory wentylacyjne, przepusty przewodów, okna wyświetlaczy.

Zasada, którą warto zapamiętać, dotyczy zależności między wymiarem szczeliny a długością fali. Szczelina zaczyna skutecznie promieniować, gdy jej najdłuższy wymiar zbliża się do połowy długości fali zakłócenia. Dla 1 GHz połowa długości fali wynosi 15 cm, dla



Fotografia 13. Obudowa z profilu aluminiowego – konstrukcja zapewniająca ciągłość ekranu na trzech ściankach przy jednoczesnym odprowadzaniu ciepła

5 GHz – 3 cm. Oznacza to, że w urządzeniach z modułami radiowymi pracującymi w pasmach gigahercowych szczelina o długości kilku centymetrów, wizualnie niepozorna, potrafi zniweczyć efekt ekranowania. Środkiem zaradczym jest skrócenie szczelin przez zwiększenie liczby punktów styku obudowy – śrub, zatrzasków, sprężystych taśm stykowych – a nie zwiększanie grubości ścianek.

Istotna jest też ciągłość galwaniczna między częściami obudowy. Powłoki lakiernicze i anodowe są izolatorami, dlatego w obudowach, w których ekranowanie ma znaczenie, przewiduje się powierzchnie styku pozbawione powłoki albo wykończone powłoką przewodzącą, na przykład chromianowaną.

W obudowach z tworzyw ekranowanie wymaga nanieśienia warstwy przewodzącej od wewnątrz – natryskiem farby z wypełniaczem metalicznym, metalizacją próżniową lub wklejką z folii. Wszystkie te metody podnoszą koszt jednostkowy i komplikują produkcję, dlatego decyzję o ekranowaniu obudowy plastikowej podejmuje się zwykle dopiero wtedy, gdy zawiodą środki układowe: filtracja na wejściach, dławiki wspólne, staranne prowadzenie mas i ekranowanie na poziomie płytki.

Kastomizacja, czyli gdzie kończy się katalog

To zjawisko, które w ostatnich latach najsilniej zmieniło rynek obudów, i jednocześnie obszar, w którym konstruktor ma najwięcej do zyskania.

Przypomnijmy liczbę przytoczoną wcześniej: usługi indywidualizacji i obróbki mechanicznej wskazało jako najważniejszy trend zmieniający rynek 68% ankietowanych dostawców obudów, a możliwość zamówienia wersji na zamówienie znalazła się wśród kryteriów wyboru dostawcy na czwartym miejscu, ze wskazaniem 52% badanych. Producenci obudów przejęli zatem znaczną część prac, które kiedyś wykonywał u siebie producent urządzenia.

Zakres usług oferowanych dziś standardowo obejmuje: wiercenie i frezowanie otworów pod gniazda, przełączniki i wyświetlacze; wykonanie i montaż paneli czołowych; nadruk grafiki i opisów, także pełnokolorowy; lakierowanie w wybranym kolorze z palety RAL; integrację klawiatur membranowych; montaż szyby ochronnej nad wyświetlaczem; montaż uszczelnień, przepustów kablowych i dławików; a coraz częściej także integrację gotowego wyświetlacza lub ekranu dotykowego wraz z pakietem usług wsparcia. Rezultatem jest obudowa dostarczana w stanie, w którym konstruktorowi pozostaje zamontować płytkę i skrócić całość.

Osobną wartością są konfiguratorzy internetowe udostępniane przez producentów systemów obudów. Pozwalają one złożyć obudowę z elementów systemu, wskazać położenie i kształt wycięć, dobrać kolor i nadruk, a następnie pobrać model 3D i wygenerować numer katalogowy wersji zamawianej. Dla konstruktora oznacza to możliwość wstawienia rzeczywistej geometrii obudowy do modelu urządzenia jeszcze przed złożeniem zamówienia.

Granica, przy której warto się zatrzymać, jest wykonanie w pełni indywidualne, wymagające własnej formy wtryskowej.

Rozwinięciem tego wątku jest druk przestrzenny. Jako technologia produkcji obudów seryjnych pozostaje mało konkurencyjny – jednostkowy koszt wydruku profesjonalnego, powiększony o obróbkę wykończeniową, przewyższa cenę wtrysku już przy niewielkich seriach. Natomiast jako narzędzie prototypowania przed wykonaniem formy jest dziś standardem: kilka wydruków pozwala zweryfikować ergonomię, dostępność złączy i montaż mechaniczny zanim zapadnie decyzja o wydatku na oprzyrządowanie.

Wzornictwo i interfejs użytkownika

Ostatnia dekada przyniosła w obudowach dla elektroniki wyraźną ewolucję wzorniczą. Prostopadłościan z widocznymi łbami wkrętów przestał być normą; producenci

NRE i MOQ – dlaczego własna forma jest droga

NRE (Non-Recurring Engineering) to bezzwrotne koszty przygotowania produkcji: projekt, uzgodnienia, prototypy, serie próbne oraz – pozycja dominująca – wykonanie oprzyrządowania, przede wszystkim formy wtryskowej. Koszt formy dla obudowy o umiarkowanej złożoności liczy się w dziesiątkach, a przy konstrukcjach wieloelementowych w setkach tysięcy złotych. Rzadko zdarza się, aby klient pokrywał go z góry w całości; częściej jest rozkładany na wiele partii, co podnosi cenę jednostkową o niewielki ułamek, ale wymaga zdefiniowania minimum zakupowego (MOQ), typowo rzędu tysiąca sztuk. Dla urządzeń specjalistycznych, produkowanych w krótkich seriach, takie minimum bywa nie do przyjęcia – i właśnie dlatego indywidualizacja wyrobów katalogowych jest dla tego segmentu rozwiązaniem lepszym niż wykonanie full custom. Drugim argumentem jest czas: projekt, poprawki, prototypy, kolejne poprawki, uzgodnienia i akceptacje potrafią przesunąć premierę rynkową o kilka kwartałów.



Fotografia 14. Klawiatury membranowe i moduły ekranów dotykowych przeznaczone do zabudowy w obudowach katalogowych

inwestowali w nowoczesne obrabiarki, wtryskarki i centra CNC, a efekty widać w geometrii wyrobów i w jakości powierzchni.

Dla konstruktora ma to konsekwencję praktyczną. Urządzenie profesjonalne konkuruje dziś nie tylko parametrami, ale i tym, jak wygląda w szafie klienta obok aparatury innych producentów. W badaniu IRE ciekawe wzornictwo i stylistykę wskazało jako istotny trend 39% ankietowanych dostawców – tyle samo, co możliwość instalacji wielu płytek i złączy oraz prostotę mocowania

komponentów. Wzornictwo przestało być kategorią oddzieloną od techniki.

Po stronie interfejsu użytkownika standardem stały się klawiatury membranowe wykonywane według projektu klienta, eliminujące płątaninę przewodów między panelem a płytką i pozwalające umieścić w jednej warstwie grafikę, opisy, okna wyświetlaczy i diody sygnalizacyjne. Coraz częściej producenci obudów oferują wersje serii katalogowych z fabrycznie zintegrowanym wyświetlaczem – od prostego modułu alfanumerycznego z klawiaturą

Lista kontrolna doboru obudowy

1. Gdzie urządzenie będzie zamontowane: na szynie w rozdzielnicach, na szynie w szafie sterowniczej, na ścianie, w wycięciu panelu, na stole, w kasce 19-calowej, na zewnątrz budynku?
2. Czy obowiązuje ograniczenie gabarytowe wynikające z normy lub z praktyki instalacyjnej – na przykład zgodność z DIN 43880 dla aparatury w rozdzielnicach?
3. Jaki stopień ochrony jest wymagany i czy dotyczy on całej obudowy, czy tylko strony czołowej? Czy konieczna jest deklaracja odporności na uderzenia w postaci kodu IK?
4. Jaka jest moc strat urządzenia i jaka jest maksymalna temperatura otoczenia w miejscu zabudowy? Czy dopuszczalne są otwory wentylacyjne?
5. Czy urządzenie jest zasilane z sieci i podlega ocenie według normy wyrobu? Jakie właściwości materiałowe obudowy będą wymagane i czy producent je deklaruje?
6. Czy wymagane jest ekranowanie elektromagnetyczne? Jeśli tak, czy obudowa metalowa nie koliduje z anteną lub z wymaganiami izolacyjnymi?
7. Ile płytek ma się zmieścić w obudowie i w jakim układzie? Jaki jest sposób ich mocowania i czy narzuca on wymiary konturu?
8. Jaka technika przyłączeniowa jest oczekiwana przez odbiorcę i czy producent obudowy oferuje gotowe pokrywy z zaciskami?
9. Czy urządzenie ma być modułowe i czy potrzebne jest złącze magistralowe na szynie?
10. Jaka jest planowana wielkość serii i horyzont czasowy? Czy indywidualizacja wersji katalogowej wystarczy, czy uzasadnione jest wykonanie własnej formy?
11. Jakie usługi wykończeniowe są potrzebne: obróbka mechaniczna, nadruk, panel czołowy, klawiatura, integracja wyświetlacza?
12. Jaki jest deklarowany termin dostawy dla wersji standardowej, a jaki dla wersji modyfikowanej?

membranową po pojemnościowy ekran dotykowy o przekątnej rzędu 2,4 cala wbudowany w pokrywę obudowy na szynę DIN, z gotowym interfejsem elektrycznym po stronie płytki. Dla producenta urządzeń automatyki budynkowej takie rozwiązanie skraca projekt o cały podzespół.

Procedura doboru

Na koniec proponujemy uporządkowaną listę pytań, którą warto przejść, zanim obudowa zostanie wpisana do zestawienia materiałowego. Kolejność nie jest przypadkowa – odpowiedzi na pytania wcześniejsze eliminują największą liczbę kandydatów.

Podsumowanie

Rynek obudów dostępny konstruktorowi w Polsce jest głęboki i – wbrew obiegu opinii – mało scentralizowany. Wyroby oferują krajowi producenci obudów metalowych i z tworzyw, oddziały producentów zagranicznych, dystrybutorzy specjalizowani, u których obudowy

stanowią istotną część biznesu, oraz dystrybutorzy katalogowi mający je w ofercie obok podzespołów. Zestawienia dostawców wraz z zakresem oferty – z podziałem na wersje metalowe, plastikowe, szafy i usługi – publikowane są w corocznym Informatorze Rynkowym Elektroniki; jest to najszybsza droga do zbudowania krótkiej listy dostawców do zapytania ofertowego.

Dwa wnioski wydają się warte podkreślenia. Pierwszy: dobór obudowy należy przeprowadzić przed projektem obwodu drukowanego, bo to obudowa narzuca więcej ograniczeń niż odwrotnie. Drugi: granica między katalogiem a wykonaniem indywidualnym przesunęła się mocno w stronę klienta i wiele problemów, które kiedyś wymagały własnej formy wtryskowej, rozwiązuje się dziś obróbką mechaniczną wyrobu katalogowego, przy zachowaniu terminu dostawy zbliżonego do standardowego. Warto z tej możliwości korzystać świadomie – pytanie do dostawcy o zakres usług indywidualizacji kosztuje jeden telefon, a potrafi zaoszczędzić kwartał pracy.

WM, KZ

Bibliografia

- Informator Rynkowy Elektroniki 2026, AVT – rozdział „Obudowy i szafy dla elektroniki i przemysłu”, strony drukowane 132..138: wyniki badania ankietowego wśród dostawców, segmentacja rynku, zestawienie ofert dostawców obecnych na rynku krajowym.
- ep.com.pl, dział „Wybór konstruktora” – wcześniejsze opracowania o obudowach i klawiaturach; obecny artykuł celowo ustawiony w opozycji do ujęcia warsztatowego, na poziomie wymagań wyrobu seryjnego.
- elektronikaB2B.pl oraz elportal.pl – serwisy własne wydawnictwa. Artykuły i prezentacje o obudowach modułowych, obudowach ze stali nierdzewnej, klawiaturach foliowych i panelach czołowych; wykorzystane jako źródło części materiału ilustracyjnego.
- Materiały katalogowe producentów obudów dostępnych na rynku krajowym (wykaz źródeł ilustracji powyżej).
- Normy przywołane w tekście: PN-EN 60529 (stopnie ochrony IP), PN-EN 62262 (kod IK), EN 60715 (szyna montażowa TS35), DIN 43880 (wymiały aparatury instalacyjnej), UL 94 (palność tworzyw), PN-EN IEC 62368-1, PN-EN 61010-1.




<https://skroc.pl/bSmnQEf>

Panel przedni najważniejszym elementem obudowy

Obudowa urządzenia elektronicznego pełni szereg niezwykle istotnych funkcji – chroni umieszczone w niej układy przed uszkodzeniami mechanicznymi oraz wpływem środowiska zewnętrznego, stanowi ramę montażową do ułożenia płytek drukowanych, podzespołów zasilania i ew. dodatkowych mechanizmów, tworzy barierę przeciwporażeniową, a nierzadko także pomaga w spełnieniu wymogów kompatybilności elektromagnetycznej, bądź bierze udział w formowaniu systemu chłodzenia. Elementem obudowy wymagającym szczególnej uwagi projektanta jest panel czołowy – w wielu przypadkach to jedyna „aktywna” część obudowy, tak ściśle zintegrowana z elektroniką i pracująca w roli elementu interfejsu użytkownika (HMI).

Zaawansowane funkcje oscyloskopów cyfrowych dla niezaawansowanych (1)

Jak dokonywać trudniejszych pomiarów i unikać pomyłek

Oscyloskopy cyfrowe od dłuższego czasu oferują szereg funkcji i ustawień pozwalających na wykonywanie wielu użytecznych pomiarów. Kiedyś wiele z tych funkcjonalności dostępnych było wyłącznie w instrumentach najwyższej klasy, obecnie jednak oferują je nawet oscyloskopy budżetowe, często bez dodatkowych opłat.

Wielu użytkowników, zwłaszcza początkujących, nie zdaje sobie sprawy ani z pełnych możliwości oscyloskopów cyfrowych, ani z ich ograniczeń. Wielu też pomija dokładną lekturę instrukcji obsługi, w której wszystkie funkcje przyrządu są szczegółowo opisane, choć nie zawsze wyczerpująco wyjaśniono w niej niektóre pułapki pomiarowe. W niniejszym artykule omówione zostaną te zagadnienia oraz zaprezentowane praktyczne wskazówki, o których instrukcje obsługi rzadko wspominają. Należy pamiętać, że nie każdy oscyloskop dysponuje kompletem funkcji, a poszczególni producenci stosują odmienne nazewnictwo i strukturę menu. Dlatego lektura dokumentacji jest zalecana nawet wtedy, gdy Czytelnik dobrze zna swój przyrząd.

Tryby i typy wyzwalania

Oscyloskopy cyfrowe, w przeciwieństwie do analogowych, oferują trzy podstawowe tryby pracy: automatyczny (Auto), normalny (Normal) oraz jednorazowy (Single). Tryb normalny odświeża przebieg za każdym razem, gdy warunek wyzwalania zostanie spełniony, prezentując również fragment sygnału poprzedzający moment wyzwolenia (pre-trigger). Aby skorzystać z tego trybu, należy wcześniej znać parametry poszukiwanego sygnału oraz właściwie dobrać typ i próg wyzwalania. Tryb automatyczny działa identycznie, lecz jeśli po upływie



określonego czasu kryterium wyzwalania nie zostanie spełnione, przebieg i tak zostanie wyświetlony. Czas ten jest zazwyczaj odwrotnością częstotliwości odświeżania oscyloskopu dla wybranej podstawy czasu. Jest to domyślny tryb pracy DSO, pozwalający na wstępne odnalezienie sygnału oraz ustalenie progu wyzwalania. W trybie Single rejestrowany i wyświetlany jest tylko jeden, pierwszy przebieg spełniający zdefiniowane kryteria. Tryb ten okazuje się nieoceniony przy badaniu stanów nieustalonych oraz nieprawidłowości w pracy układów, na przykład zjawisk towarzyszących włączaniu i wyłączaniu zasilania. W połączeniu z zaawansowanymi typami wyzwalania pozwala w łatwy sposób przechwycić zdarzenia rzadkie i sporadyczne. Ułatwia również pracę w sytuacjach, w których operator musi w pełni skupić się na precyzyjnym przyłożeniu sondy do punktu pomiarowego.

Większość oscyloskopów, poza wskazaniem jednego z kanałów analogowych jako źródła wyzwalania, udostępnia wejście wyzwalania zewnętrznego (Ext) lub możliwość synchronizacji sygnałem sieci zasilającej (AC Line). Ta ostatnia opcja jest szczególnie przydatna podczas projektowania i serwisowania przetwornic impulsowych z aktywnym układem PFC bądź falowników synchronizowanych z siecią

elektroenergetyczną, stosowanych w instalacjach fotowoltaicznych. Zewnętrzny kanał wyzwalania przydaje się wtedy, gdy oscyloskop współpracuje z generatorem funkcyjnym przy wielokanałowych pomiarach obwodów analogowych, uwalniając kanały pomiarowe do rejestracji sygnałów. Warto pamiętać, że wejściowy tor wyzwalania ma własne, niezależne od torów pionowych filtry. W oscyloskopie Siglent SDS1104X-U ustawienia te są powiązane ze sprzężeniem i obejmują wybory: AC, DC, LF Reject oraz HF Reject. Sprzężenia AC i DC działają analogicznie jak w kanałach pomiarowych. Filtr LF Reject tłumi składowe od DC do częstotliwości rzędu kilkuset kiloherców lub pojedynczych megaherców, co zapobiega przedwczesnemu wyzwaniu przez przydźwięk lub szumy niskoczęstotliwościowe. Z kolei HF Reject odcina składowe powyżej 1,2 MHz, eliminując ryzyko fałszywego wyzwania na szpilkach zakłócających, pochodzących od przetwornic impulsowych czy szybkich układów cyfrowych. Dodatkowo w menu wyzwalania dostępna jest funkcja Noise Rejection, która zwiększa histerezę komparatorów, podnosząc odporność na szum kosztem nieznacznego spadku czułości. Istotnym parametrem jest również czas blokady wyzwalania (Trigger Holdoff). Określa on minimalny czas, jaki musi upłynąć od ostatniego wyzwolenia do uaktywnienia możliwości kolejnego wyzwolenia. Pozwala to ustabilizować obraz złożonych przebiegów powtarzalnych, na przykład synchronizując wyzwalanie zawsze na początku ramki w protokołach szeregowych pozbawionych sprzętowej obsługi w danym przyrządzie, takich jak 1-Wire firmy Dallas/Maxim.

Podstawowym typem wyzwalania jest detekcja zbocza (Edge) – narastającego, opadającego lub obu naraz. Nie wyczerpuje to jednak możliwości współczesnych przyrządów. Standardem stało się wyzwalanie szerokością impulsu (Pulse Trigger). Definiuje się w nim polaryzację oraz kryterium czasowe: czas trwania impulsu mniejszy, większy, mieszczący się w zadanym oknie czasowym lub niemieszczący się w nim. Typ ten sprawdza się doskonale przy weryfikacji sygnałów PWM oraz w diagnostyce niestandardowych transmisji szeregowych. Przykładowo, pozwala sprawdzić, czy układ podrzędny (Slave) magistrali 1-Wire po odebraniu impulsu zerującego (Reset Pulse, stan niski trwający powyżej 480 μ s) odczeka właściwy czas 15...60 μ s, zanim odpowie impulsem obecności (Presence Pulse) o czasie trwania 60...240 μ s. Pokrewnym rozwiązaniem jest wyzwalanie interwałem czasowym (Interval), gdzie warunkiem jest czas upływający pomiędzy dwoma kolejnymi zboczami tej samej polaryzacji. Umożliwia to wykrywanie nieregularności w przerwaniach transmisji czy niestabilności częstotliwości sygnałów zegarowych. Wyzwalanie zanikiem sygnału (Dropout) następuje z kolei wtedy, gdy po zarejestrowaniu zbocza kolejne nie pojawi

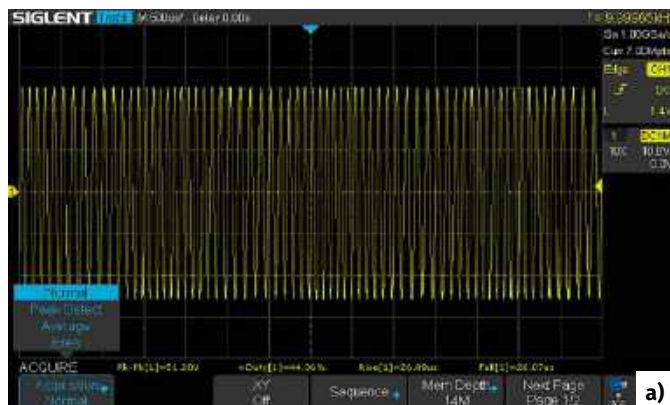
się w zdefiniowanym czasie, co sygnalizuje niepożądane zawieszenie magistrali. Obserwując na drugim kanale linię zasilania, można natychmiast zweryfikować, czy zanik transmisji nie wynika ze spadku napięcia. Do badania stanów przejściowych zasilania służy wyzwalanie czasem narastania lub opadania (Slope Trigger). W trybie Single pozwala ono skontrolować odpowiedź zasilacza przy włączeniu pod kątem prawidłowego działania układów POR (Power-On Reset) i BOR (Brown-Out Reset) w mikrokontrolerach. Polega ono na zdefiniowaniu dwóch progów napięciowych oraz dopuszczalnego czasu przejścia sygnału między nimi, co pozwala także oceniać szybkość przełączania stromych zboczy w magistralach cyfrowych.

Wiele oscyloskopów oferuje nadal wyzwalanie analogowym sygnałem wideo, co współcześnie znajduje zastosowanie głównie w pracach ze sprzętem retro, systemami monitoringu CCTV starego typu czy układami OSD w modułarstwie FPV.

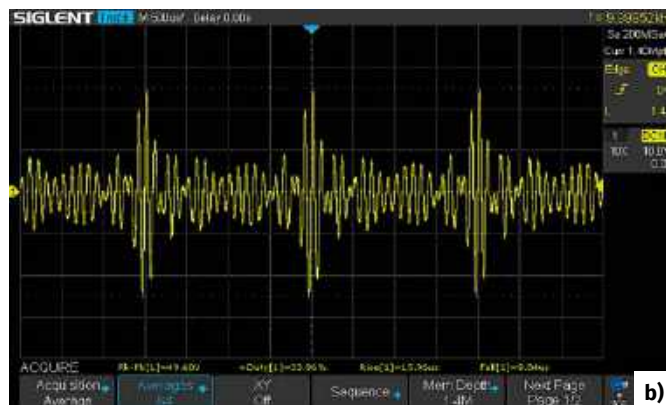
Wyzwalanie okienkowe (Window Trigger) jest niezastąpione podczas projektowania i diagnozowania analogowych torów audio. Wykorzystuje dwa niezależne progi napięciowe: dolny i górny. W trybie bezwzględny (Absolute) definiuje się konkretne wartości obu progów, natomiast w trybie względnym (Relative) ustala się poziom odniesienia oraz symetryczną tolerancję. Wyzwalanie następuje przy przekroczeniu okna przez sygnał, co ułatwia wykrywanie przesterowań wzmacniaczy lub sprawdzanie poziomów logicznych w układach cyfrowych. Pokrewnym typem jest wyzwalanie impulsem niepełnym (Runt). Układ wyzwala podstawę czasu, gdy sygnał przekroczy pierwszy próg napięciowy, ale zawróci przed osiągnięciem drugiego. Jest to kluczowe narzędzie do wychwytywania stanów metastabilnych, zakłóceń szyn danych oraz niepełnych przełączeń poziomów logicznych.

Niezwykle użyteczne są również typy wyzwalania logicznego (Pattern) oraz magistralami szeregowymi (Serial). Tryb logiczny analizuje kombinację stanów logicznych na wejściach analogowych (zdefiniowanych na podstawie indywidualnych progów napięciowych) w powiązaniu z funkcjami boolowskimi AND, OR, NAND lub NOR. Wyzwalanie magistralami szeregowymi pozwala z kolei na synchronizację z konkretnymi zdarzeniami w transmisjach I²C, SPI, UART/RS-232, CAN czy LIN (takimi jak bit startu, bajt adresu, brak potwierdzenia NACK czy określona wartość danych). Większość przyrządów umożliwia równoległe dekodowanie protokołów w czasie rzeczywistym i prezentację danych w oknie tabeli zdarzeń. Dawniej pakiety dekodujące wymagały kosztownych licencji programowych, obecnie standardem rynkowym staje się ich bezpłatne fabryczne odblokowanie.

W nowoczesnych oscyloskopach cyfrowych podstawowe mechanizmy wyzwalania realizowane są sprzętowo



a)



b)

Rysunek 1. Sygnał 10 kHz z modulacją częstotliwości FM w trybie akwizycji normalnej (a) oraz zniekształcenie powstałe w wyniku zastosowania trybu uśredniania (b)

za pomocą szybkich komparatorów i przetworników DAC połączonych ze strukturami logicznymi FPGA. W instrumentach budżetowych bardziej złożone wyzwalanie protokołami szeregowymi bywa dzielone pomiędzy logikę FPGA (detekcja bitów startu i stopu) a oprogramowanie procesora głównego (wyszukiwanie rozbudowanych sekwencji danych), co przy dużych prędkościach transmisji może ograniczać tempo akwizycji.

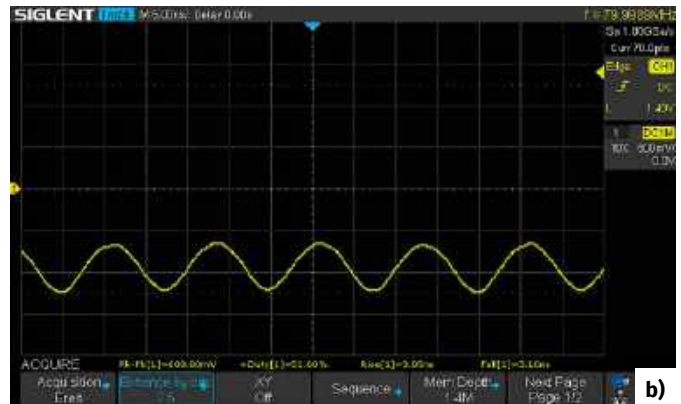
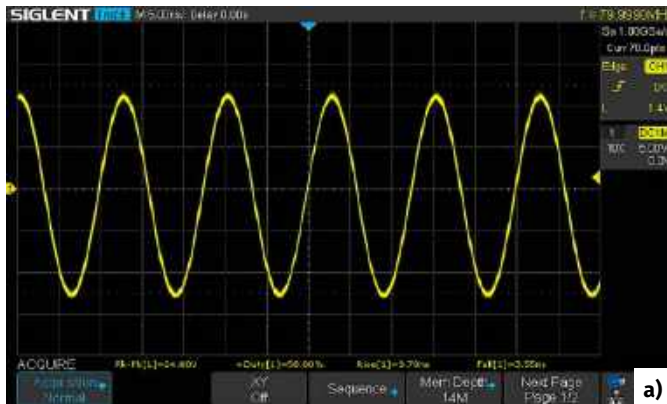
System akwizycji i jego możliwości

Zadaniem systemu akwizycji jest próbkowanie sygnałów z przetworników ADC, zapisywanie próbek w szybkiej pamięci oraz ich dalsze przetwarzanie przez dedykowany układ ASIC lub FPGA w celu wyświetlenia na ekranie. Kluczowymi parametrami są: pasmo analogowe, częstotliwość próbkowania w czasie rzeczywistym oraz głębokość pamięci akwizycji. Przykładowo oscyloskop Siglent SDS1104X-U oferuje pamięć 14 Mpts i maksymalne próbkowanie 1 GSa/s, co dla pojedynczego kanału pozwala zarejestrować odcinek czasu o długości 14 ms z pełną rozdzielczością próbek. Należy mieć na uwadze, że w oscyloskopach czterokanałowych tej klasy przetworniki ADC oraz pamięć są współdzielone – przy aktywnych czterech kanałach próbkowanie spada do 250 MSa/s, a pamięć do 3,5 Mpts na kanał. Zgodnie z twierdzeniem Nyquista-Shannona do wiernego odtworzenia sygnału częstotliwość próbkowania musi być ponad dwukrotnie wyższa od najwyższej składowej harmonicznej. W praktyce pomiarowej przyjmuje się współczynnik bezpieczeństwa wynoszący co najmniej 2,5 dla sygnałów sinusoidalnych oraz minimum 4...5 dla sygnałów impulsowych. Przy próbkowaniu 250 MSa/s oscyloskop o paśmie 100 MHz pracuje na granicy poprawnej rekonstrukcji przebiegów, co grozi zjawiskiem aliasingu. Do odtwarzania ciągłego kształtu fali ze skwantowanych próbek stosuje się standardowo interpolację $\sin(x)/x$, natomiast interpolacja liniowa łączy próbki odcinkami prostymi, co lepiej oddaje strome zbocza sygnałów cyfrowych, lecz zniekształca przebiegi sinusoidalne.

W menu akwizycji użytkownik ma do wyboru różne tryby pracy. Tryb normalny (Normal) pobiera próbki w regularnych odstępach czasu. Tryb detekcji szczytów (Peak Detect) zapisuje wartości minimalne i maksymalne w każdym okresie próbkowania, tworząc obwiednię sygnału i umożliwiając rejestrację bardzo wąskich szpilek i zakłóceń, które w trybie normalnym zostałyby pominięte wskutek decymacji. Tryb uśredniania (Average) redukuje losowy szum nieskorelowany poprzez uśrednianie kolejnych przebiegów wyzwalanych synchronicznie, uwypuklając powtarzalne składowe sygnału. Należy jednak pamiętać o pułapce: jeśli mierzony sygnał cechuje się niestabilnością fazową (jitter) lub modulacją częstotliwości, uśrednianie całkowicie zniekształci wyświetlany obraz, co przedstawiono na **rysunku 1**. Odrębną funkcją jest tryb podwyższonej rozdzielczości pionowej High-Res (w oscyloskopach Siglent: Eres). Wykorzystuje on filtrację cyfrową nadpróbkowanego sygnału (oversampling), zwiększając rozdzielczość przetwornika ADC z 8 do nawet 11 bitów i skutecznie eliminując szum, jednak kosztem drastycznego zawężenia pasma przepustowego. Zmniejszenie pasma dla każdego dodatkowego bitu rozdzielczości sprawia, że szybkie sygnały zostają silnie stłumione, co demonstruje **rysunek 2**, gdzie amplituda sygnału 80 MHz uległa niemal całkowitej redukcji po włączeniu filtra Eres.

Domyślnym trybem prezentacji danych jest tryb YT, przedstawiający napięcie w funkcji czasu. Alternatywny tryb XY przypisuje napięcie jednego kanału do osi pionowej, a drugiego do osi poziomej, co pozwala na generowanie figur Lissajous. Stanowi to klasyczną metodę pomiaru przesunięcia fazowego w obwodach analogowych:

- Generator funkcyjny podłącza się do badanego obwodu, doprowadzając sygnał wejściowy do kanału pierwszego, a sygnał z wyjścia obwodu do kanału drugiego.
- W trybie XY dokonuje się odczytu wartości przecięć osi i maksymalnych wychyleń elipsy zgodnie z **rysunkiem 3**.



Rysunek 2. Sygnał 80 MHz z generatora funkcyjnego: przebieg w trybie akwizycji normalnej o amplitudzie $2,4 V_{pp}$ (a) oraz silne tłumienie amplitudy do wartości 60 mV po włączeniu trybu Eres wskutek cyfrowego ograniczenia pasma przepustowego (b)

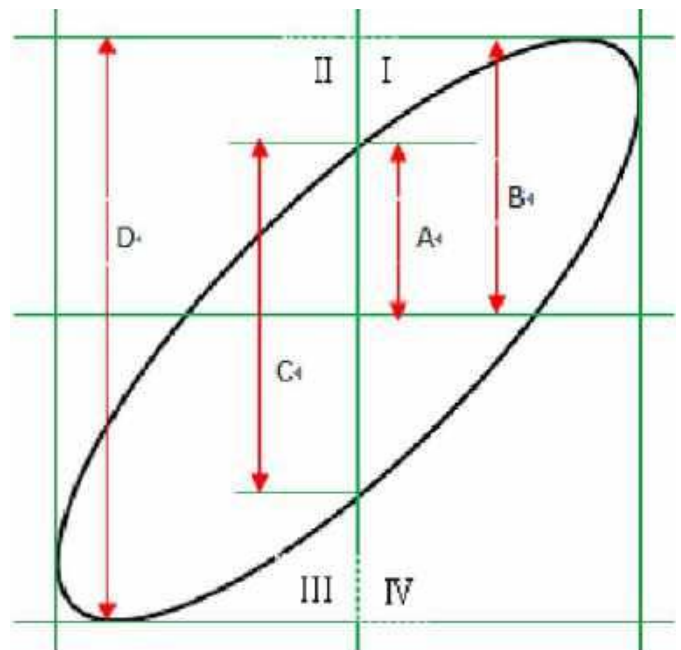
- Przesunięcie fazowe wyznacza się z zależności:

$$\theta = \arcsin\left(\frac{A}{B}\right) = \arcsin\left(\frac{C}{D}\right)$$

Choć obecnie pomiary fazy realizowane są automatycznie za pomocą kursorów lub wbudowanych funkcji kreślenia charakterystyk częstotliwościowych (wykresów Bodego), tryb XY pozostaje użytecznym narzędziem do natychmiastowej oceny symetrii i przesunięć w torach sygnałowych.

Wiele przyrządów wyposażono w wyjście pomocnicze (Aux Out), które można skonfigurować jako wyjście wyzwalania (Trigger Out) lub wyjście testu maski (Pass/Fail). Test granic maski jest powszechnie stosowany w kontroli produkcyjnej do automatycznej weryfikacji zgodności kształtu sygnału z zadaniem wzorcem tolerancji. Z kolei wyjście Trigger Out w warunkach laboratoryjnych umożliwia precyzyjną synchronizację zewnętrznych generatorów funkcyjnych, mierników częstotliwości czy analizatorów widma, realizując sprzężone pomiary odpowiedzi impulsowej układów.

Wyjątkowo efektywną funkcją jest pamięć segmentowa (nazywana w przyrządach Siglent trybem sekwencyjnym). Pozwala ona na podzielenie całkowitej pamięci akwizycji na określoną liczbę segmentów i rejestrowanie wyłącznie ramek zawierających zdarzenie wyzwalające, z pominięciem bezczynnych odcinków czasu. Przykładem zastosowania jest magistrala I²C pracująca z prędkością 400 kbit/s, gdzie pakiety o czasie trwania 112,5 μs przesyłane są co 2 ms. W tradycyjnym trybie akwizycji ciągła rejestracja kilkuset milisekund wymusza kompromis pomiędzy długością zapisu a szybkością próbkowania, przez co ponad 90% pamięci marnowane jest na rejestrację stanu spoczynkowego szyny. W trybie segmentowym oscyloskop rejestruje wyłącznie pakiety danych przy zachowaniu maksymalnej częstotliwości próbkowania, co pozwala zgromadzić w pamięci serię rzadkich błędów (np. ramek z bitem NACK) w celu ich późniejszego porównania, korelacji ze stanem linii zasilających



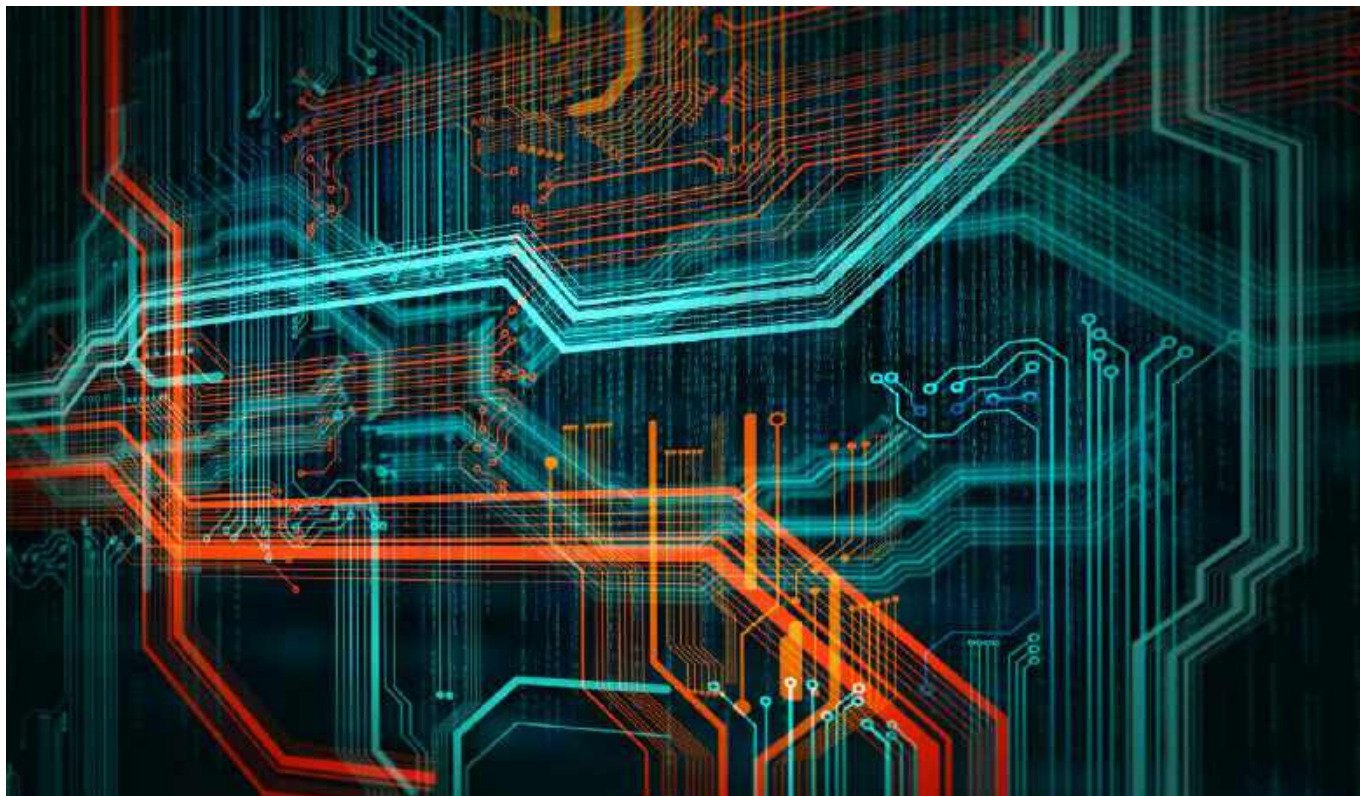
Rysunek 3. Pomiar przesunięcia fazowego w trybie XY. Dla sygnałów w fazie (0°) obraz ma postać linii prostej o nachyleniu dodatnim, wraz ze wzrostem przesunięcia przekształca się w elipsę (dla 90° w elipsę symetryczną lub okrąg), by w przeciwfazie (180°) stać się linią prostą o nachyleniu ujemnym. Wyznaczenie odcinków A i B lub C i D pozwala na precyzyjne obliczenie kąta fazowego

oraz automatycznego zdekodowania w tabeli zdarzeń. Połączenie pamięci segmentowej z detekcją szczytów Peak Detect umożliwia bezbłędną lokalizację sporadycznych zaburzeń elektromagnetycznych w systemach wbudowanych.

Podsumowanie

W niniejszym artykule omówiono wybrane zaawansowane możliwości torów wyzwalania i systemów akwizycji współczesnych oscyloskopów cyfrowych. Właściwe wykorzystanie tych narzędzi pozwala w pełni wykorzystać potencjał posiadanego przyrządu pomiarowego, skracając czas diagnozy i eliminując podstawowe błędy interpretacyjne.

Paweł Kowalczyk



Projekt płytki na zlecenie

W tym numerze rozpoczynamy w rubryce Repetytorium cykl o płytkach drukowanych – od projektu, przez wytwarzanie, po montaż. Cykl pokazuje, ile decyzji podejmuje konstruktor, zanim pakiet danych trafi do wytwórcy, i jak każda z nich odkłada się na koszcie gotowego wyrobu. Naturalnym uzupełnieniem tego obrazu jest pytanie, które prędzej czy później pojawia się w każdej pracowni: czy tę pracę wykonać samodzielnie, czy zlecić. Firma PCBWay, znana polskim konstruktorom przede wszystkim jako wytwórca płytek prototypowych, prowadzi od kilku lat również dział projektowy. Poniżej przedstawiamy zakres tej usługi, sposób zamawiania i to, co warto wiedzieć przed rozmową o wycenie.

Trzy usługi, nie jedna

Oferta jest podzielona na trzy linie, które można zamawiać osobno albo łącznie. Rozróżnienie ma znaczenie praktyczne, bo od niego zależy, co konstruktor musi dostarczyć i co dostanie w zamian.

Layout płytki

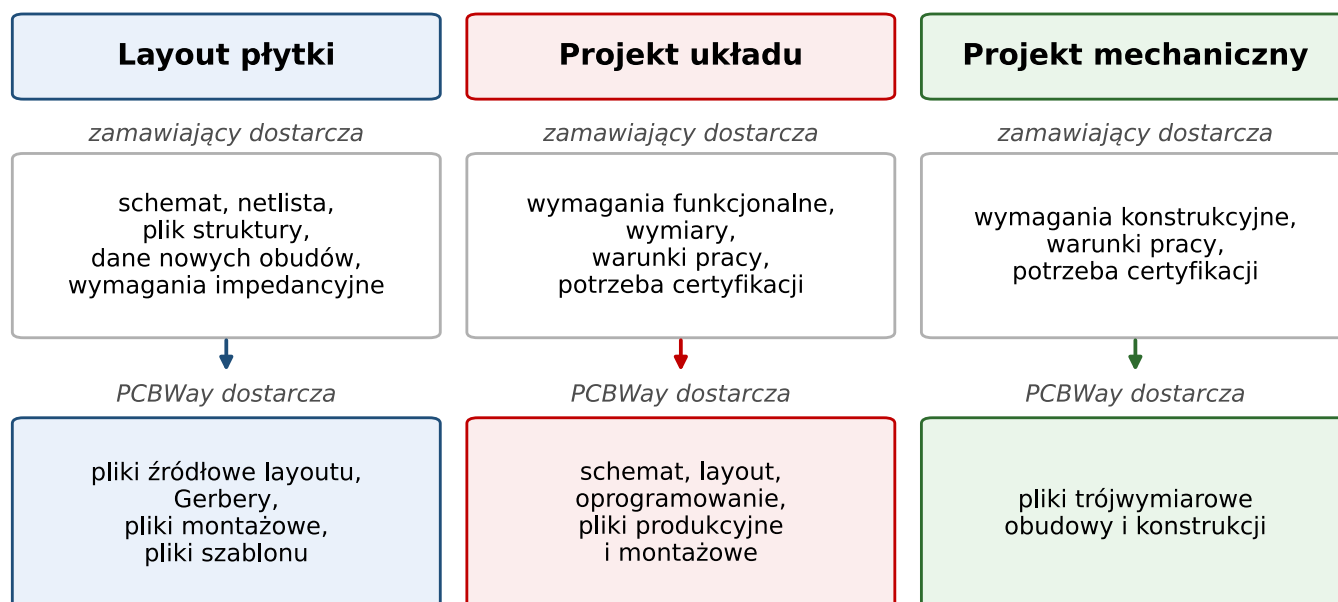
Wariant dla konstruktora, który ma gotowy schemat i chce zlecić samo rozmieszczenie i trasowanie. PCBWay deklaruje obsługę projektów do stu tysięcy wyprowadzeń i płytek do sześćdziesięciu warstw, w tym konstrukcji sztywnych, elastycznych, sztywno-elastycznych, HDI, wysokoczęstotliwościowych i mocy. W zakres wchodzi projekt stosu warstw oraz kontrolowana impedancja.

Konstruktor dostarcza schemat, netlistę, plik struktury mechanicznej, dane elementów, dla których trzeba utworzyć nowe pola lutownicze, oraz wymagania dotyczące impedancji i prądów. Na wyjściu otrzymuje pliki źródłowe layoutu, Gerbery, pliki montażowe i pliki szablonu.

Projekt układu

Wariant szerszy: od wymagań funkcjonalnych do działającego egzemplarza. Obejmuje projekt elektroniczny, schemat, layout oraz oprogramowanie wbudowane. PCBWay podaje doświadczenie w rozwoju firmware'u dla mikrokontrolerów rodzin STM32, ESP32, Nordic i NXP oraz dla układów SoC.

Tu konstruktor dostarcza nie dokumentację, lecz wymagania: opis funkcji, ograniczenia wymiarowe, warunki pracy wyrobu oraz informację, czy urządzenie ma



Model współpracy decyduje o tym, czy zamawiający otrzymuje pliki źródłowe — patrz rys. 3

Rysunek 1. Trzy linie usług projektowych PCBWay: co dostarcza zamawiający i co otrzymuje. Opracowanie redakcji na podstawie materiałów PCBWay – pcbway.com/pcbdesign.html

podlegać certyfikacji. Na wyjściu otrzymuje pliki produkcyjne, montażowe i szablonu, a w zależności od wybranego modelu współpracy także pliki źródłowe.

Projekt mechaniczny i obudowa

Trzecia linia obejmuje konstrukcję mechaniczną i obudowę, z plikami trójwymiarowymi jako produktem wyjściowym. W praktyce ma to znaczenie tam, gdzie płytka i obudowa powstają równolegle i muszą do siebie pasować – czyli w większości wyrobów przeznaczonych na rynek, a nie do szafy sterowniczej.

Decyzje projektowe a koszt wytwarzania

Najciekawszy fragment materiałów technicznych PCBWay dotyczy tego, co dla czytelnika tego pisma jest sednem sprawy: w jaki sposób decyzje podjęte przy ekranie przekładają się na rachunek za produkcję. Firma opisuje to wprost i zachęca zamawiających, by przy składaniu zlecenia podawali docelowy przedział kosztu wyrobu – po to, żeby zespół projektowy mógł dobrać rozwiązania pod ten budżet, zamiast optymalizować je po fakcie.

Wskazane czynniki są następujące.

- Liczba warstw. Każda kolejna warstwa to dodatkowe operacje procesowe, dłuższy czas produkcji i wyższe wymagania technologiczne. PCBWay deklaruje dążenie do minimalizacji liczby warstw przy zachowaniu wymagań kompatybilności elektromagnetycznej i integralności sygnału, z pierwszeństwem dla konstrukcji dwuwarstwowych tam, gdzie są wystarczające.
- Materiał. Standardowy laminat FR-4 jest najkorzystniejszy kosztowo. Materiały wysokociepłotliwościowe,

o podwyższonej temperaturze zeszklenia albo przeznaczone do sygnałów szybkich stosuje się wtedy, gdy są rzeczywiście potrzebne.

- Rodzaj i średnica przelotek. Otwory na wylot są najtańsze. Przelotki wiercone laserowo oraz ślepe i zagrzebane wymagają dodatkowych operacji i wydłużają produkcję.
- Szerokość ścieżki i odstęp. Węższe ścieżki i ciaśniejsze odstępki wymagają precyzyjniejszego wyposażenia i podnoszą odsetek wyrobów wadliwych, a przez to koszt.
- Wymiary i kształt płytki. Wielkość i kontur decydują o wykorzystaniu materiału. Kształty prostokątne są korzystniejsze kosztowo, przy czym wymiar musi pomieścić wszystkie elementy i pasować do obudowy.
- Dobór elementów. Obudowy nietypowe wymagają precyzyjniejszego montażu. Elementy o drobnym rastrze wymuszają trudniejszy projekt. Osobną sprawą jest dostępność: elementy mające zamienniki stabilizują koszt, jednoźródłowe wystawiają wyrób na wahań cen i opóźnień.

Wszystko to jest zgodne z tym, co opisujemy w części pierwszej i trzeciej naszego cyklu, i warto zwrócić uwagę na jedną rzecz: to są argumenty wytwórcy, który jednocześnie projektuje. Zespół projektowy pracujący w firmie posiadającej własne fabryki zna swoje tabele możliwości technologicznych z pierwszej ręki i ma bezpośredni interes w tym, żeby projekt dał się wykonać za pierwszym razem.

Firma podkreśla przy tym granicę, o której łatwo zapomnieć w rozmowie o oszczędnościach: projektowanie wyłącznie pod kątem obniżenia kosztu prowadzi

Liczba warstw	każda warstwa to dodatkowe operacje procesowe <i>w materiale źródłowym wskazana jako jeden z najistotniejszych czynników</i>
Materiał laminatu	FR-4 najtańszy; wysoka Tg i materiały RF drożej
Rodzaj przelotek	otwory na wylot najtańsze; ślepe i laserowe drożej
Szerokość ścieżki i odstęp	węższe wymagają precyzji i podnoszą braki
Wymiary i kształt	prostokąt lepiej wykorzystuje panel
Obudowy elementów	drobny raster wymusza trudniejszy montaż
Dostępność elementów	zamienniki stabilizują koszt; jedno źródło ryzykowne

Kolejność na rysunku odpowiada kolejności w materiale źródłowym PCBWay i nie jest uszeregowaniem według ważności — wymienione czynniki stanowią listę elementów wpływających na koszt wytwarzania.

Rysunek 2. Czynniki projektowe wpływające na koszt wytwarzania, w kolejności podanej w materiale źródłowym. Opracowanie redakcji na podstawie materiałów technicznych PCBWay

do wyrobu, który nie spełnia wymagań funkcjonalnych. Optymalizacja ma polegać na znalezieniu równowagi, a nie na cięciu.

Ile kosztuje projekt i dlaczego tyle

PCBWay odnosi się wprost do zarzutu, który słyszy od zamawiających: że koszt projektu wydaje się wysoki w zestawieniu z kosztem produkcji. Argumentacja firmy sprowadza się do trzech punktów.

Po pierwsze, projekt płytki to nie rysowanie schematu. To praca obejmująca reguły elektryczne, kontrolowaną impedancję, kompatybilność elektromagnetyczną, zarządzanie ciepłem i niezawodność, wykonywana przez zespół łączący kompetencje sprzętowe, programistyczne i pomiarowe. Proces jest iteracyjny, a przy płytkach szybkich, wielowarstwowych, HDI oraz przy wyrobach podlegających wymaganiom motoryzacyjnym albo medycznym – wyraźnie dłuższy i obciążony większym ryzykiem.

Po drugie, pojedynczy błąd projektowy potrafi kosztować wielokrotność honorarium za projekt. Stąd nacisk na trafienie za pierwszym razem i na ograniczenie liczby serii prototypowych.

Po trzecie – i to jest argument, który w rozmowie z konstruktorem brzmi najlepiej, bo jest sprawdzalny we własnym doświadczeniu – firma przytacza przypadki zamawiających, którzy dla oszczędności wybrali tańszy zespół projektowy i zapłacili za trzy nieudane serie prototypowe. Suma czasu i kosztów przekroczyła cenę projektu wykonanego profesjonalnie.

Warto policzyć samodzielnie

Rachunek, który proponujemy przeprowadzić przed rozmową o zleceniu, jest prosty. Po jednej stronie: honorarium za projekt. Po drugiej: koszt jednej dodatkowej serii prototypowej – płytki, elementy, montaż, wysyłka, cło – pomnożony przez spodziewaną liczbę iteracji, plus własny czas na diagnostykę i poprawki, wyceniony według rzeczywistej stawki godzinowej, plus koszt opóźnienia wdrożenia. Jeżeli druga strona rachunku wychodzi wyższa, decyzja jest arytmetyczna, a nie ambicjonalna.

Trzy modele współpracy

Element oferty, który wyróżnia ją na tle typowych biur projektowych i który wymaga świadomej decyzji, to trzy warianty rozliczenia różniące się tym, co zamawiający otrzymuje na własność.

To jest miejsce, w którym trzeba czytać uważnie. Model trzeci jest najtańszy, ale wiąże zamawiającego z jednym wytwórcą na cały cykl życia wyrobu. Model pierwszy kosztuje najwięcej i daje pełną swobodę, łącznie z możliwością przeniesienia produkcji gdzie indziej albo wprowadzania zmian własnymi siłami. Wybór nie jest kwestią budżetu na projekt, tylko decyzją o tym, gdzie ma leżeć własność dokumentacji wyrobu – a to jest decyzja strategiczna, nie kosztowa. Zastrzeżenie, które trzeba postawić wyraźnie, żeby nie było nieporozumienia: różnica między modelami dotyczy tego, jakie pliki zamawiający otrzymuje, a nie tego, do kogo należą prawa do projektu.

Prawa te we wszystkich trzech modelach pozostają przy zamawiającym. PCBWay deklaruje, że chroni własność intelektualną klientów i że nie wykorzystuje ani nie odsprzedaży powierzonych projektów, a wszystkie przekazane informacje i pliki traktuje jako

Model 1 cena: najwyższa pliki źródłowe pliki produkcyjne uruchomiony egzemplarz	Model 2 cena: średnia pliki produkcyjne uruchomiony egzemplarz	Model 3 cena: najniższa uruchomiony egzemplarz
pełna swoboda: zmiany i produkcja gdziekolwiek	produkcja dowolnie, ale zmiany projektowe wracają do PCBWay	produkcja seryjna związana z PCBWay

We wszystkich modelach: bezpłatne wsparcie techniczne, stały rabat 5% na montaż w PCBWay po zakończeniu projektu, zniżka na kolejne wersje projektu

We wszystkich modelach prawa do projektu pozostają przy zamawiającym.
Brak plików źródłowych nie oznacza przeniesienia praw — PCBWay deklaruje, że nie wykorzystuje ani nie odsprzedaży projektów klientów, a przekazane materiały są poufne.

Rysunek 3. Trzy modele współpracy i ich konsekwencje dla dostępu do dokumentacji wyrobu. Prawa do projektu we wszystkich modelach pozostają przy zamawiającym. Opracowanie redakcji na podstawie cennika usług i deklaracji PCBWay

Model	Co otrzymuje zamawiający	Cena	Uwagi
Model 1	Pliki źródłowe projektu, pliki produkcyjne oraz uruchomiony egzemplarz	najwyższa	Pełna dokumentacja; zamawiający może samodzielnie wprowadzać zmiany i produkować gdziekolwiek
Model 2	Pliki produkcyjne oraz uruchomiony egzemplarz	średnia	Bez plików źródłowych; produkcja możliwa u dowolnego wytwórcy, ale zmiany projektowe wracają do PCBWay
Model 3	Uruchomiony egzemplarz	najniższa	Bez plików; produkcja seryjna musi odbywać się w PCBWay

REKLAMA

Profesjonalne usługi projektowania układu PCB



Już od
50 USD
Za projekt

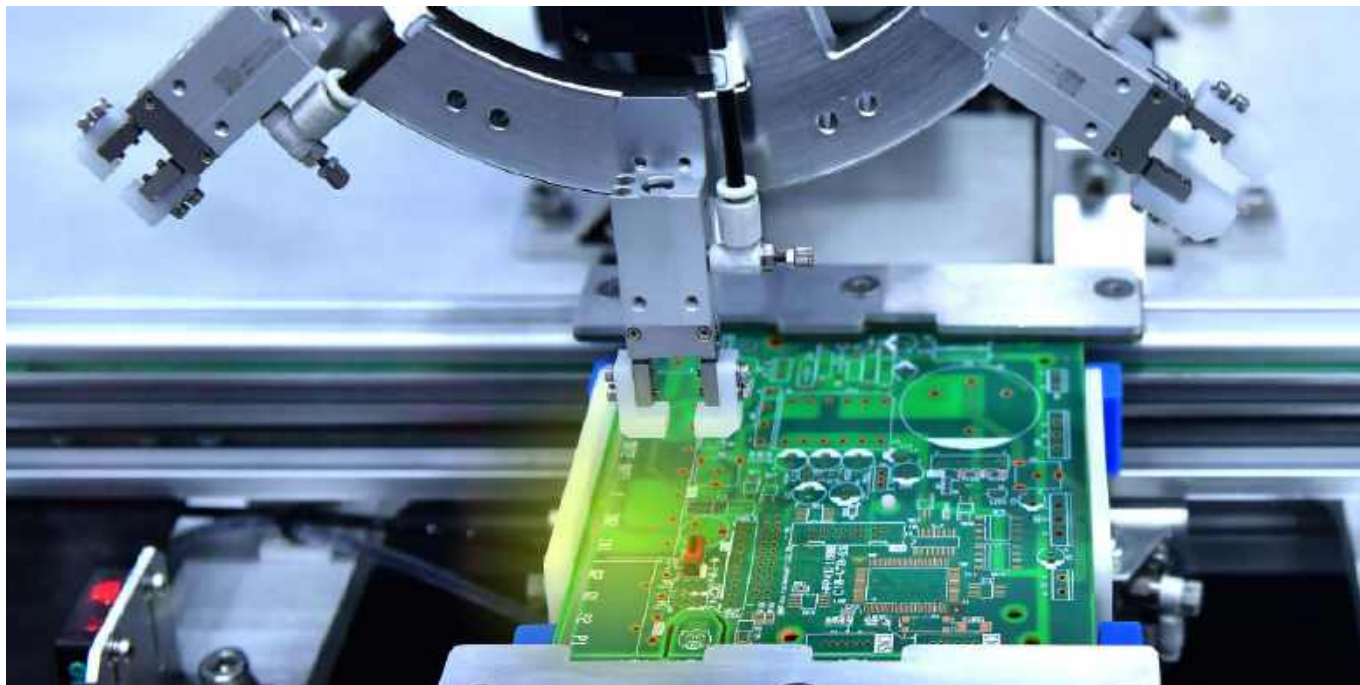
PCB do
60
Warstw

- Obsługa Altium, KiCad i innych programów
- Doświadczenie w projektowaniu PCB High-Speed, HDI i RF



 www.pcbway.com

 service@pcbway.com



poufne. Brak plików źródłowych w modelu trzecim oznacza więc ograniczenie praktyczne – zamawiający nie może samodzielnie wprowadzać zmian ani przenieść produkcji seryjnej gdzie indziej – a nie przeniesienie praw autorskich do rozwiązania. Zakres poufności i praw warto mimo to zapisać w umowie, tak jak przy każdym zleceniu projektowym.

Wszystkie trzy modele obejmują, według deklaracji firmy, bezpłatne wsparcie techniczne, stały pięcioprocentowy rabat na montaż w PCBWay po zakończeniu projektu oraz zniżkę na kolejne wersje projektu.

Jak wygląda zlecenie

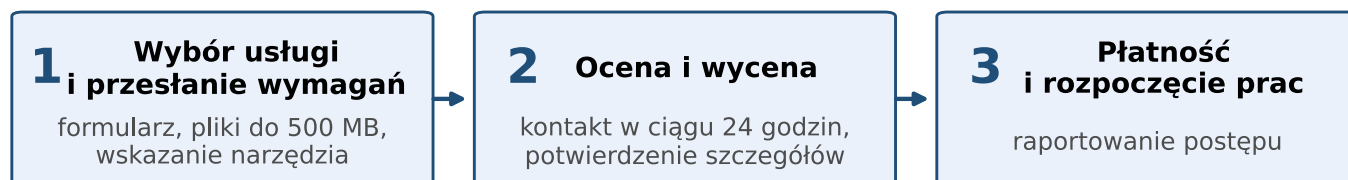
Przebieg jest trzyetapowy i w całości prowadzony przez formularz internetowy.

- Krok pierwszy: wybór usługi – layout, projekt układu albo projekt mechaniczny – i przesłanie wymagań

wraz z plikami. Formularz przyjmuje pliki do pięciuset megabajtów i pozwala wskazać narzędzie projektowe: Altium Designer, KiCad albo inne.

- Krok drugi: ocena projektu przez inżynierów, potwierdzenie szczegółów technicznych i wycena. Firma deklaruje kontakt w ciągu dwudziestu czterech godzin od zgłoszenia.
- Krok trzeci: płatność i rozpoczęcie prac, z bieżącym raportowaniem postępu.

Punkt, który warto odnotować i który odróżnia dobrze poprowadzone zlecenie od źle poprowadzonego: proces przewiduje przeglądy po stronie zamawiającego. Przy layoutcie PCBWay przekazuje do zatwierdzenia rozmieszczenie i trasowanie, a zamawiający ma potwierdzić zasadność rozmieszczenia, stosu warstw i kontroli impedancji. Przy projekcie układu do zatwierdzenia idą kolejno schemat, layout i oprogramowanie.



Punkty zatwierdzenia po stronie zamawiającego

Layout:
rozmieszczenie, trasowanie,
stos warstw, impedancja

Projekt układu:
schemat, layout,
oprogramowanie

Rysunek 4. Przebieg zlecenia z zaznaczonymi punktami zatwierdzenia po stronie zamawiającego. Opracowanie redakcji na podstawie opisu procesu PCBWay

Innymi słowy: zlecenie projektu nie zwalnia konstruktora z rozumienia tego, co dostaje. Przeglądy są miejscem, w którym wiedza z naszego cyklu przydaje się najbardziej – bo pozwala zadać właściwe pytania, zanim projekt trafi do produkcji.

Zaplecze i zakres branżowy

PCBWay podaje ponad dziesięć lat doświadczenia w tym obszarze, ponad dwudziestu inżynierów w zespole projektowym i ponad tysiąc zrealizowanych projektów. Wśród obsługiwanych branż wymienia elektronikę powszechnego użytku, automatykę przemysłową, elektronikę medyczną, motoryzację, internet rzeczy oraz energetykę.

Wyróżnikiem, na który firma kładzie nacisk, jest integracja pionowa: własne zakłady produkcji płytek, montażu, obróbki skrawaniem i wtrysku tworzyw. Z punktu widzenia zamawiającego oznacza to jednego partnera od projektu do serii oraz brak etapu przekazywania dokumentacji między firmami – a to jest etap, na którym w praktyce najczęściej gubią się szczegóły.

Firma deklaruje również współpracę z jednostkami certyfikującymi i możliwość przeprowadzenia certyfikacji CE, FCC, CSA oraz UL. Dla wyrobów przeznaczonych na rynek europejski jest to element wart uwagi na etapie planowania, ponieważ wymagania certyfikacyjne wpływają na projekt płytki, a nie tylko na dokumentację – o czym piszemy

szerzej w drugiej części cyklu, poświęconej kompatybilności elektromagnetycznej.

Dla kogo

Zamykając: usługa ma sens w trzech sytuacjach, które w polskiej praktyce konstruktorskiej występują najczęściej.

Pierwsza to projekt wykraczający poza kompetencje zespołu – płytka szybka, wielowarstwowa, z kontrolowaną impedancją albo z układem w obudowie BGA, przy braku doświadczenia z takimi konstrukcjami. Druga to brak czasu, nie brak umiejętności: projekt trzeba zrobić, a zespół jest zajęty czym innym, a termin wdrożenia nie czeka. Trzecia to sytuacja, w której zamawiający potrzebuje jednego partnera od wymagań do wyrobu, bez zarządzania łańcuchem podwykonawców.

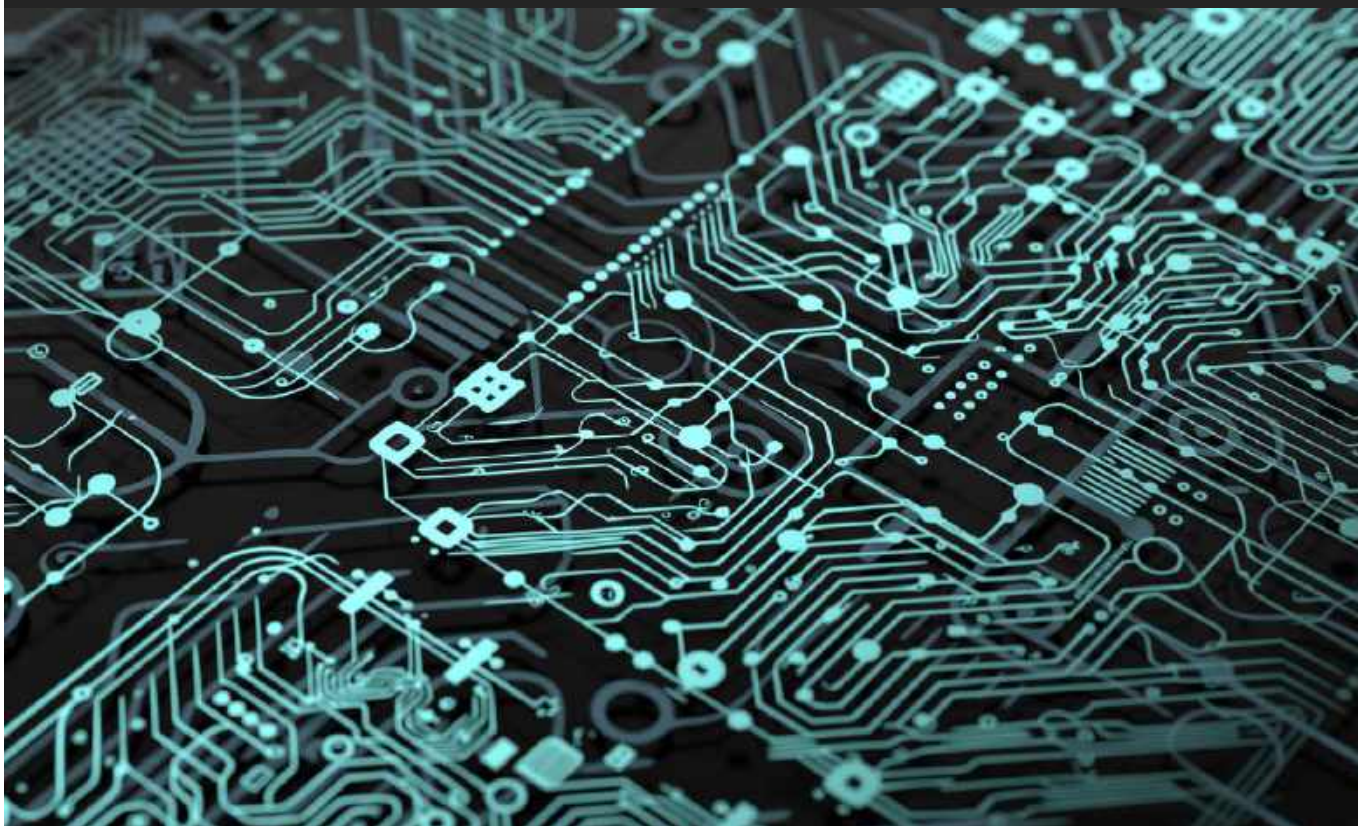
W każdym z tych przypadków punktem wyjścia jest ten sam rachunek: honorarium za projekt wobec kosztu iteracji, które przy samodzielnej pracy trzeba przewidzieć. Rachunek ten warto zrobić na kartce przed wysłaniem zapytania.

Zapytanie ofertowe i opis usług: pcbway.com/pcbdesign.html. Kontakt w sprawie współpracy: anson@pcbway.com. Materiały techniczne cytowane w tym artykule są dostępne na blogu firmy w dziale PCB Design & Layout.



Nota redakcyjna: Niniejszy materiał jest artykułem sponsorowanym, przygotowanym na zamówienie firmy PCBWay i oznaczonym zgodnie z prawem prasowym. Dane techniczne, parametry usług i deklaracje dotyczące doświadczenia pochodzą z materiałów firmowych i stron internetowych PCBWay; redakcja ich nie weryfikowała. Cykl „PCB w praktyce”, publikowany w rubryce Repetytorium, powstaje niezależnie od tej współpracy i nie podlega uzgodnieniom z reklamodawcą.

Praktyczny przewodnik konstruktora po płytkach drukowanych – od projektu przez wytwarzanie do montażu. Cykl czterech artykułów publikowanych sukcesywnie w czterech kolejnych wydaniach EP



PCB w praktyce, część 1

Projektowanie PCB

Ten cykl nie nauczy nikogo projektowania, wytwarzania i montażu płytek drukowanych. Cztery artykuły nie zastąpią podręcznika ani praktyki. Cykl „PCB w praktyce” ma inny cel: pokazać, gdzie w procesie od pomysłu do zmontowanej płytki leżą punkty decyzyjne, jakie są realne możliwości kooperacji, ile to kosztuje i czego trzeba się nauczyć, żeby dalej radzić sobie samodzielnie. Część pierwsza dotyczy projektowania PCB. Tematy kolejnych trzech części: Integralność sygnału i EMC, Wytwarzanie PCB, Montaż modułów na PCB.

1. Co naprawdę decyduje o powodzeniu projektu

Konstruktor, który po raz pierwszy zamawia płytkę w Shenzhen, zwykle spodziewa się, że najtrudniejsze będzie trasowanie. Po kilku projektach wie już, że trasowanie jest tą częścią pracy, która idzie najgładziej. Kłopoty biorą się gdzie indziej: z pola lutowicznego przepisane z niewłaściwej strony noty katalogowej, z płaszczyzny masy przeciętej ścieżką sygnałową, z pliku wierceń wyeksportowanego w calach, gdy reszta pakietu jest w milimetrach, z opisu montażowego zachodzącego na pole lutownicze na czarnej masce.

Wspólna cecha tych błędów polega na tym, że każdy z nich powstaje w kilkanaście sekund, a wykrywa się go tygodnie później i w zupełnie innym miejscu procesu. To jest realna oś całego cyklu – nie „jak zaprojektować płytkę”, lecz „gdzie leżą punkty, w których tania decyzja powoduje drogie skutki” (tabela 1).

Warto zwrócić uwagę, czego w tej tabeli nie ma. Nie ma mnożników kosztu w rodzaju „1 – 10 – 100 – 1000”, które od lat krążą po prezentacjach branżowych. Nie udało się ustalić badania, z którego miałyby pochodzić, więc nie podajemy ich jako faktu. Kierunek zależności jest zresztą oczywisty i w zupełności wystarcza: koszt naprawy

Tabela 1.

Etap wykrycia błędu	Co trzeba powtórzyć	Czego nie da się odzyskać
Podczas rysowania schematu	fragment schematu	nic
Podczas trasowania	fragment layoutu	nic
Po kontroli DRC i przeglądzie DFM	eksport danych, ponowne zgłoszenie	pół dnia pracy
Po odebraniu gotych płytek	całe zamówienie płytek	koszt płytek, czas produkcji i transportu
Po montażu	płytki i część elementów	Koszt płytek i montażu, elementy w obudowach nienadających się do przelutowania
Po uruchomieniu u odbiorcy	wszystko powyższe	zaufanie odbiorcy i termin wdrożenia

rośnie skokowo przy każdym przekroczeniu kolejnej granicy między etapami, a najkosztowniejsza granica przebiega między gotową płytką a zmontowanym urządzeniem.

Z tego wynika jedna zasada praktyczna, do której będziemy wracać w każdym odcinku. Pracę, którą można wykonać wcześniej, należy wykonać wcześniej, nawet gdy wydaje się nudna. Sprawdzenie pola lutowniczego w nocie katalogowej producenta jest nudne. Przepisanie tabeli możliwości technologicznych wytwórcy do reguł projektowych jest nudne. Obejrzenie własnych Gerberów w niezależnej przeglądarce jest nudne. Każda z tych czynności zwraca się przy pierwszym unikniętym błędzie. Uwaga – opisując oferty rynkowe podajemy często ceny; są to informacje orientacyjne, gdyż mogą się z czasem zmieniać.

2. Warsztat: czym się dziś projektuje i ile to kosztuje

Wybór narzędzia bywa przedstawiany jako kwestia gustu. W praktyce jest to decyzja o trzech konsekwencjach: ile zapłacisz rocznie, czy odzyskasz swoje dane, gdy przestaniesz płacić, oraz komu będziesz mógł przekazać projekt do dalszej pracy. Wygoda pracy jest dopiero czwarta na tej liście.

Rok 2026 przyniósł w tej sprawie zdarzenie, które warto omówić na samym początku, ponieważ dotyczy niewielkiej grupy polskich konstruktorów.

2.1. Koniec EAGLE

7 czerwca 2026 roku Autodesk zakończył sprzedaż i wsparcie programu EAGLE. Jako dalszą drogę producent wskazuje subskrypcję Autodesk Fusion i przestrzeń roboczą Electronics, która obejmuje do 999 arkuszy schematu, 16 warstw i nieograniczoną powierzchnię płytki. Subskrypcje miesięczne i roczne EAGLE przestały być dostępne i przestało być możliwe uruchomienie programu oraz praca w nim. Zmiana nie objęła licencji wieczystych.

To jest istotniejsze, niż wygląda na pierwszy rzut oka. Program nie przestał być rozwijany – on przestał się uruchamiać, ponieważ wyłączono serwery licencyjne. Konstruktor, który ma w archiwum projekty sprzed pięciu lat, ma je w formacie, do którego nie istnieje już czynny edytor, o ile nie dysponuje starą licencją wieczystą.

Wnioski są dwa. Pierwszy dotyczy tych, którzy jeszcze mają dane w plikach z rozszerzeniem BRD: należy je przenieść teraz, do Fusion albo do KiCada, i zarchiwizować równolegle w postaci neutralnej, czyli w Gerberach z atrybutami. Drugi jest ogólniejszy i wróci w tym artykule jeszcze kilkakrotnie – format zamknięty ma wartość dopóty, dopóki jego właściciel uważa, że mu się to opłaca.

Bezpłatny wariant EAGLE nie istnieje. Autodesk kieruje użytkowników do Fusion for Personal Use, które obejmuje zintegrowane CAD, CAM i CAE, dwa arkusze schematu, dwie warstwy sygnałowe i ograniczoną powierzchnię płytki. Komunikat producenta wraz z odpowiedziami na pytania: autodesk.com/support/technical/article/caas/sfdcarticles/sfdcarticles/Autodesk-EAGLE-Announcement-Next-steps-and-FAQ.html

2.2. Narzędzia przemysłowe

Do tej grupy zaliczamy programy przeznaczone do projektów seryjnych, wielowarstwowych i szybkich, z rozbudowanym zarządzaniem ograniczeniami projektowymi i z zapleczem analitycznym.

Altium Designer

Najczęściej spotykany w polskich biurach konstrukcyjnych program z tej półki. Model licencyjny zmienił się i to jest pierwsza rzecz, którą trzeba zrozumieć przed rozmową z handlowcem. Sprzedaż idzie dziś przez pakiet Altium Develop, zbudowany wokół wspólnej przestrzeni roboczej. Do pięciu miejsc typu Author, przy czym każde kolejne kosztuje 995 dolarów rocznie. Przy więcej niż pięciu autorach ECAD producent kieruje do wyższego pakietu Altium Agile Teams. CircuitStudio nie jest już sprzedawany nowym klientom, a dotychczasowi użytkownicy zachowują dostęp do produktu i dokumentacji.

W pakiecie startowym są: jedno miejsce autorskie, jedna przestrzeń robocza na wspólne projekty i biblioteki oraz nieograniczona liczba użytkowników przestrzeni z prawem do przeglądania, komentowania, zarządzania zadaniami i dostępu do listy materiałowej, ale bez prawa edycji. Z funkcji objętych pakietem podstawowym warto odnotować projektowanie hierarchiczne i wielokanałowe, projektowanie sterowane regułami, natywne

środowisko trójwymiarowe, projektowanie wiązek oraz zestawy wielopłytkowe.

Cennik i opis pakietu: altium.com/develop/pricing

Cadence: OrCAD X i Allegro X

Rodzina o szerszej rozpiętości niż oferta Altium – od samego edytora schematów OrCAD X Capture, przez OrCAD X Standard i OrCAD X Professional, po Allegro X dla projektów najbardziej wymagających oraz Allegro X AI ze wspomaganiem trasowania. Ceny są publikowane w sklepie internetowym dystrybutora, co w tym segmencie zdarza się rzadko i warto z tego skorzystać przy negocjacjach.

OrCAD X Standard kosztuje od 1440 do 5790 euro, zależnie od rodzaju licencji. Warianty są trzy. Licencja roczna imienna jest przypisana do konkretnego użytkownika i konkretnej maszyny. Licencja roczna pływająca jest obsługiwana przez serwer licencji i może być współdzielona w obrębie jednej lokalizacji, przy czym w danej chwili korzysta z niej jeden użytkownik. Licencja wieczysta pływająca daje dostęp na dziewięćdziesiąt dziewięć lat, ale wymaga rocznej umowy serwisowej, aby po pierwszym roku dalej otrzymywać aktualizacje.

Zakres OrCAD X Standard obejmuje wstępną analizę integralności sygnału, trasowanie par różnicowych, kontrolę reguł w czasie rzeczywistym, menedżer ograniczeń, współpracę ECAD-MCAD w trzech wymiarach, analizę wykonalności oraz – co ważne przy migracji – translatory z Altium, EAGLE i PADS.

Sklep i cennik: ema-eda.com/store/pcb-products/orcad-x-standard/

Siemens EDA: Xpedition i PADS Professional

Xpedition jest narzędziem korporacyjnym wycenianym indywidualnie, obecnym przede wszystkim w motoryzacji, lotnictwie i telekomunikacji. PADS Professional plasuje się niżej i jest kierowany do mniejszych zespołów konstrukcyjnych, z zachowaniem menedżera ograniczeń i analizy wykonalności. Ceny obu produktów Siemens EDA podaje na zapytanie. Opis obu rodzin: eda.sw.siemens.com/en-US/pcb/

Zuken: CR-8000 i eCADSTAR

W Polsce spotykane rzadziej, ale obecne w motoryzacji i wszędzie tam, gdzie projekt obejmuje rozbudowaną wiązkę elektryczną. Ceny na zapytanie. Warto natomiast odnotować ruch ze stycznia 2026 roku: Zuken udostępnił wspólnie z firmą LogicSwap Solutions bezpłatny moduł migracyjny pozwalający użytkownikom Cadence OrCAD przenieść projekty PCB, schematy i biblioteki do eCADSTAR. Moduł jest bezpłatny dla licencjonowanych użytkowników eCADSTAR i pobiera się go ze strony LogicSwap.

Dla konstruktora rozważającego zmianę narzędzia to sygnał wart odnotowania: bariera wyjścia z jednego ekosystemu komercyjnego została obniżona celowo.

2.3. Narzędzia dla małych zespołów i pojedynczych konstruktorów

KiCad

Nie jest to już „darmowa alternatywa”. Wersja 10 ukazała się 20 marca 2026 roku i powstała z 7609 unikatowych zmian w kodzie i tłumaczeniach – zauważalnie więcej niż przy wersji 9. Do bibliotek dodano 952 nowe symbole, 1216 pól lutowniczych i 386 modeli trójwymiarowych. Ponad 78 procent pól lutowniczych w bibliotece jest dziś generowanych z danych, a nie rysowanych ręcznie, co ma bezpośrednie przełożenie na powtarzalność. Domyślnym formatem modeli trójwymiarowych jest STEP i od wersji 10 wyłącznie STEP, co zmniejszyło rozmiar instalacji i poprawiło zgodność geometrii z eksportem.

Z nowości wersji 10 najistotniejsze dla codziennej pracy są cztery.

- Kompletna przebudowa strojenia długości ścieżek, obejmująca możliwość definiowania ograniczeń w dziedzinie czasu zamiast wyłącznie w jednostkach długości, oraz profile strojenia pozwalające zdefiniować parametry trasowania osobno dla każdej warstwy. Dla par różnicowych i magistrali taktowanych jest to różnica jakościowa, nie kosmetyczna.
- Graficzny edytor reguł projektowych, w pełni zgodny z dotychczasowym językiem reguł niestandardowych. Można zacząć od definicji graficznej i płynnie przejść do zapisu tekstowego, gdy reguła robi się złożona.
- Warianty projektu, czyli równoległe wersje jednego przedsięwzięcia wspólne co do schematu, ale różniące się właściwościami elementów, na przykład listą materiałową.
- Importery projektów z Allegro, PADS oraz gEDA i Lepton PCB.

Ostatnia pozycja ma wymowę handlową, nie tylko techniczną: droga z narzędzia komercyjnego do otwartego przestała być jednokierunkowa.

Z pozostałych zmian warto wymienić tryb ciemny w systemie Windows, konfigurowalne paski narzędzi, cofanie zmian wewnątrz okien dialogowych, zaznaczanie lassem, obiekty na warstwach wewnętrznych w polach lutowniczych, zamianę wyprowadzeń i sekcji z propagacją między schematem a płytką, bloki projektowe w edytorze płytki, obsługę kodów kreskowych i eksport trójwymiarowego PDF.

Materiały: kicad.org, dokumentacja docs.kicad.org, konwencje bibliotek klc.kicad.org. Konferencja KiCon Europe odbywa się 7–9 września 2026 roku.

Autodesk Fusion Electronics

Naturalna droga dla tych, którzy pracują równolegle w mechanice, i jedyna droga dla dawnych użytkowników EAGLE, którzy chcą zostać u tego samego producenta. Ograniczenie warte rozważenia jest natury organizacyjnej, nie technicznej: środowisko działa w modelu chmurowym, co w części firm bywa problemem regulaminowym.

Pozostałe

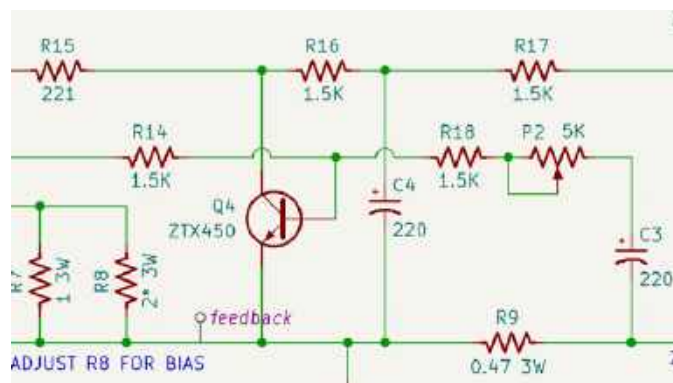
DipTrace, Proteus PCB Design, DesignSpark PCB i Target 3001! mają w mniejszych pracowniach ustaloną pozycję, każdy z nieco innym profilem: Proteus wyróżnia się symulacją współbieżną z modelem mikrokontrolera, DesignSpark jest bezpłatny i powiązany z dystrybutorem, Target 3001! ma silną pozycję w krajach niemieckojęzycznych. Aktualne cenniki i wersje próbne wszystkich czterech są dostępne na stronach producentów: diptrace.com, labcenter.com, rs-online.com/designspark, ibfriedrich.com

Na trzecim poziomie są narzędzia przeglądarkowe i otwarte. EasyEDA w wersjach Standard i Pro jest powiązany kapitałowo i procesowo z JLCPCB oraz LCSC, co jest zaletą i wadą zarazem: integracja z magazynem części i wyceną produkcji nie ma sobie równych, ale projekt staje się trudniejszy do przeniesienia gdzie indziej. LibrePCB i Horizon EDA to projekty otwarte o mniejszej skali niż KiCad, warte uwagi ze względu na czytelne, dobrze udokumentowane modele danych.

2.4. Jak podejść do wyboru

Cztery pytania, w tej właśnie kolejności.

- Czy odzyskam swoje dane po zakończeniu subskrypcji? Przypadek EAGLE pokazuje, co znaczy odpowiedź przecząca. Niezależnie od narzędzia archiwizuj każde wydanie projektu również w Gerberach z atrybutami – to jedyny format, który przetrwa zmianę dostawcy.
- Komu przekażę projekt? Jeżeli współpracujesz z biurem konstrukcyjnym albo zlecasz layout na zewnątrz,



Rysunek 1. Przecięcia typu hop-over na schemacie – konwencja czytelności wprowadzona w KiCadzie 10. Ilustracja źródłowa: kicad.org/img/blog/2026/version-10-release/hopovers.png (KiCad, licencja CC BY 3.0)

narzędzie wybiera się razem z partnerem, a nie przed jego wyborem.

- Czy potrzebuję prawdziwego menedżera ograniczeń? Przy projektach z kontrolowaną impedancją i dopasowaniem długości różnica między narzędziem z menedżerem ograniczeń a narzędziem z listą reguł jest różnicą kilku dni pracy na projekt.
- Dopiero teraz: ile to kosztuje? Przy stawce godzinowej konstruktora różnica między KiCadem a pakietem komercyjnym zwraca się lub nie zwraca w kilkunastu godzinach oszczędzonej pracy rocznie. Warto ten rachunek zrobić na kartce, zamiast rozstrzygać spór światopoglądowo.

3. Schemat jest dokumentem produkcyjnym, nie rysunkiem połączeń

Schemat czyta co najmniej pięć osób: ty za rok, serwisant, technolog przygotowujący montaż, konstruktor przejmujący temat i zaopatrzeniowiec zamawiający elementy. Rysunek, który rozumie wyłącznie autor w dniu narysowania, jest kosztem odłożonym w czasie.

Konwencje, od których nie warto odstępować: zasilanie u góry, masa u dołu, przepływ sygnału od lewej do prawej. Bloki funkcjonalne na osobnych arkuszach, powiązane hierarchią, zamiast jednego arkusza formatu A1 z dwustu elementami. Oznaczenia referencyjne zgodne z IEC 60617 i przyjętą w firmie listą, a nie z doraźną inwencją.

Każdy element musi mieć w polach opisowych producenta, numer katalogowy producenta, obudowę, tolerancję oraz parametr krytyczny: napięcie pracy dla kondensatora, moc dla rezystora, prąd dla diodki. Bez tego lista materiałowa jest bezwartościowa, a przy montażu zleconym generuje serię pytań i opóźnienie, którego nikt nie planował. Do tego zagadnienia wrócimy szczegółowo w części czwartej, przy omawianiu listy materiałowej i pliku rozmieszczenia.

Kontrola ERC nie jest formalnością do odklikania. Typy wyprowadzeń przypisane w symbolach powodują, że narzędzie wykryje niepodłączone wejście, konflikt dwóch wyjść na jednej sieci albo zasilanie bez źródła. Sieć nazwana przez pomyłkę tak samo w dwóch miejscach schematu to błąd, którego nie znajdzie żadna kontrola na etapie płytki, bo na płytce jest już poprawnie zrealizowanym zwarciem.

Warto na schemacie zostawiać notatki: skąd wzięła się wartość rezystora w pętli sprzężenia, jakie założenie przyjęto co do prądu obciążenia, z której noty katalogowej i z której jej strony pochodzi układ aplikacyjny. Schemat bez notatek to schemat, którego nie da się zweryfikować bez powtórzenia całej pracy projektowej.

Punkty pomiarowe planuje się na schemacie, nie na layoutcie. Jeżeli wiadomo, że przy uruchamianiu trzeba będzie zmierzyć napięcie na wyjściu przetwornicy i zobaczyć

przebieg na linii zegara, to na schemacie pojawiają się odpowiednie punkty kontrolne z oznaczeniami, a na płytce znajdują one swoje miejsce automatycznie. Odwrotna kolejność kończy się lutowaniem drutu do wyprowadzenia obudowy QFN.

4. Biblioteki i pola lutownicze: miejsce, gdzie lenistwo zawsze się mści

Jeżeli trzeba wskazać jedną czynność, na której konstruktorzy tracą najwięcej płytek, będzie to pole lutownicze przepisane z niewłaściwego źródła.

Źródła są trzy i mają jasną hierarchię ważności. Pierwszeństwo ma zalecane pole lutownicze z noty katalogowej producenta elementu, ponieważ to producent odpowiada za lutowność swojej obudowy i to on wykonał badania. Drugie w kolejności jest pole wygenerowane zgodnie z normą IPC, przydatne wtedy, gdy producent nic nie zaleca albo gdy potrzebna jest spójność w obrębie całej biblioteki. Ostatnie miejsce zajmuje gotowe pole z biblioteki narzędzia lub z serwisu internetowego – nie dlatego, że jest z zasady złe, lecz dlatego, że nie wiadomo, kto i na jakiej podstawie je narysował.

Stan normalizacji wymaga jednego wyjaśnienia, bo bywa źródłem nieporozumień. Obowiązują IPC-7351B oraz wydany w maju 2023 roku IPC-7352, zatytułowany Generic Guideline for Land Pattern Design. Rewizja IPC-7351C, której nazwa od lat krąży w ustawieniach narzędzi, w materiałach szkoleniowych i na slajdach konferencyjnych, nigdy się nie ukazała. Jeżeli narzędzie oferuje tryb „IPC-7351C”, oznacza to tylko tyle, że jego autor przyjął pewne założenia co do dokumentu, który nie powstał.

Normy IPC opisują trzy poziomy gęstości pola lutowniczego. Poziom o największej ilości materiału daje najdłuższe pola i największy zapas na spoinę. Poziom pośredni jest przeznaczony do typowych zastosowań. Poziom o najmniejszej ilości materiału pozwala uzyskać najgęstsze rozmieszczenie, ale zakłada opanowany proces montażowy i pozostawia najmniejszą tolerancję na błąd.

Reguła praktyczna, która sprawdza się przy pracy z wytwórcami azjatyckimi i przy montażu własnym: przy

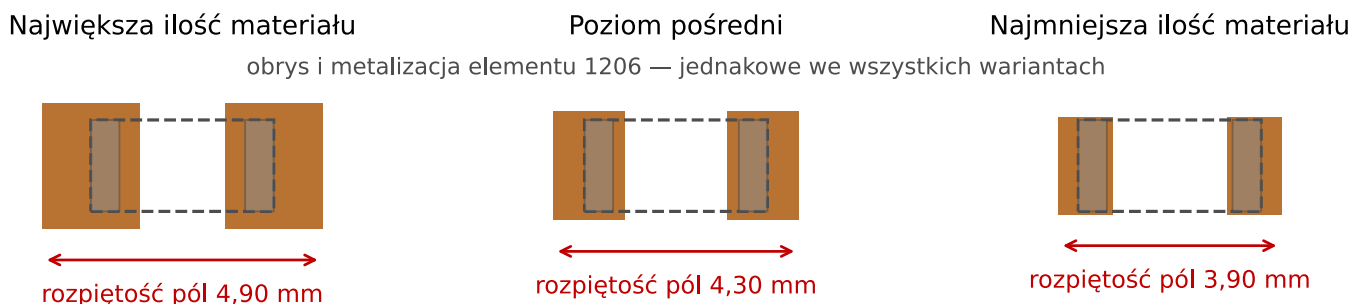
montażu zleconym wybieramy poziom pośredni, przy montażu ręcznym poziom o największej ilości materiału, ponieważ dłuższe pola upraszczają lutowanie i inspekcję, a poziom najgęstszy rezerwujemy dla płytek, na których naprawę brakuje miejsca, i tylko wtedy, gdy proces montażowy jest znany.

Pole lutownicze musi być skonfrontowane z możliwościami wytwórcy, a nie tylko z notą katalogową. Weźmy typowe wartości jednego z popularnych wytwórców azjatyckich: minimalne pole SMD ma wymiary 0,25 na 0,25 mm, minimalny odstęp między polami SMD należącymi do różnych sieci wynosi 0,15 mm, a minimalny odstęp pola od ścieżki 0,1 mm, przy czym lokalnie dla obudów BGA dopuszcza się 0,09 mm na płytkach wielowarstwowych.

Osobnej uwagi wymaga maska lutownicza, ponieważ tu kryje się pułapka, która regularnie psuje projekty z gęstymi obudowami. Minimalny mostek maski między polami wynosi 0,10 mm dla maski zielonej, czerwonej, żółtej, niebieskiej i fioletowej, ale 0,13 mm dla czarnej i białej, przy miedzi 1 oz. Przy miedzi 2 oz jest to 0,20 mm niezależnie od koloru. Innymi słowy: ta sama płytka z układem QFN o rastrze 0,4 mm może wyjść poprawnie w kolorze zielonym i nie wyjść w czarnym. Warto to wiedzieć przed wyborem koloru pod obudowę urządzenia.

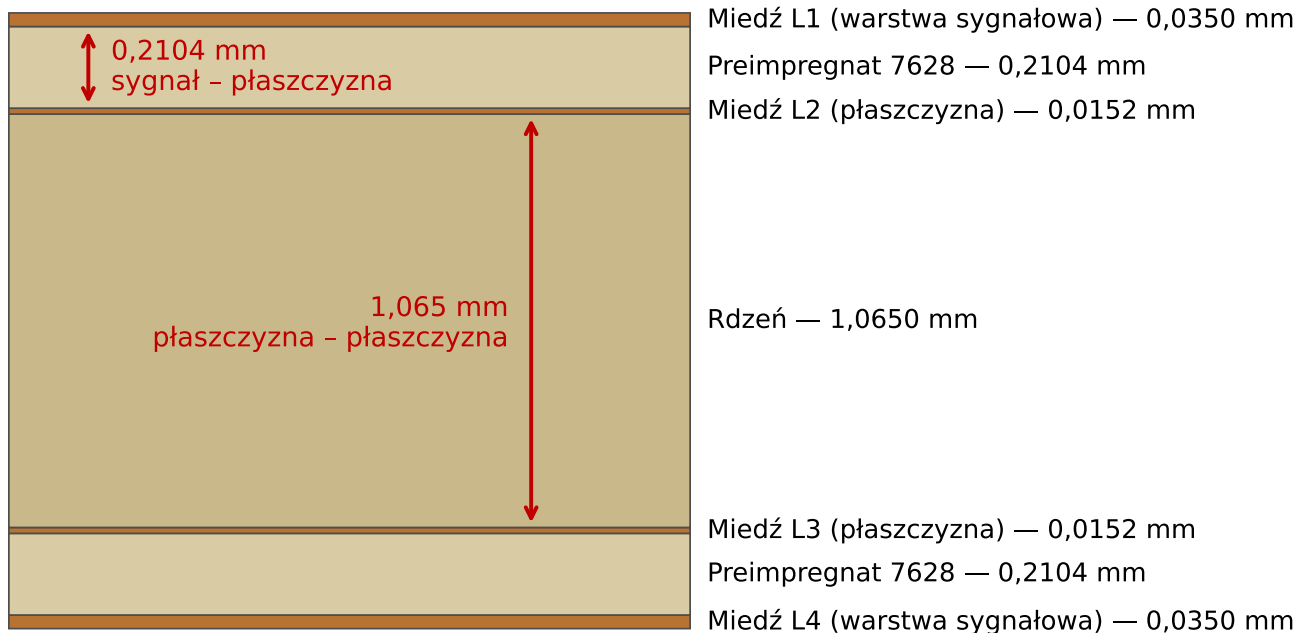
Odsłona maski nad polem jest u tego samego wytwórcy prowadzona w stosunku jeden do jednego od czerwca 2025 roku, po wymianie naświetlarek bezpośrednich, przy zaleceniu zachowania co najmniej 0,09 mm odstępu między odsłonami maski a sąsiednimi ścieżkami. Kto ma w bibliotece pola z dawnym rozszerzeniem maski, powinien to sprawdzić.

Opis montażowy podlega własnym ograniczeniom: minimalna szerokość linii 0,15 mm, minimalna wysokość znaku 1,0 mm, minimalna odległość opisu od pola lutowniczego 0,15 mm. Znaki mniejsze nie będą czytelne, a opis nachodzący na pole zostanie przez wytwórcę obcięty – albo, w gorszym wariancie, nie zostanie i pojawi się farba na powierzchni lutowniczej.



Rysunek 2. Trzy poziomy gęstości pola lutowniczego dla tego samego elementu 1206: przy niezmiennym obrysie i metalizacji elementu zmienia się zapas na spoinę i rozpiętość pól. Poziomy gęstości definiuje IPC-7352 (electronics.org)

Stos czterowarstwowy 1,6 mm — grubości rzeczywiste, skala zachowana



Rysunek 3. Przekrój stosu czterowarstwowego o grubości 1,6 mm w skali, z rzeczywistymi grubościami warstw. Widoczna dysproporcja: 0,2104 mm od sygnału do płaszczyzny wobec 1,065 mm między płaszczyznami. Rysunek własny; dane za tabelą stosów wytwórcy – jlcpcb.com/impedance

Dwie zasady organizacyjne na koniec. Po pierwsze, model trójwymiarowy elementu jest obowiązkowy, a nie ozdobny – kolizja z obudową urządzenia wykryta na ekranie kosztuje minutę, a wykryta po odebraniu płytek kosztuje serię. Po drugie, biblioteka jest częścią projektu i zmienia się razem z nim. Biblioteki „globalne”, współdzielone między projektami i modyfikowane w trakcie, sprawiają, że po dwóch latach nie da się odtworzyć płytki, którą się wtedy wyprodukowało.

Metoda weryfikacji, prosta i zaskakująco skuteczna: wydrukować pole lutownicze w skali jeden do jednego i przyłożyć do niego rzeczywisty element. Zajmuje to trzy minuty i wyłapuje pomyłkę w jednostkach, obrót o dziewięćdziesiąt stopni i przestawiony raster.

5. Stos warstw: łączyć cztery zamiast dwóch

Argument kosztowy przeciwko płytkom czterowarstwowym w warunkach roku 2026 przestał istnieć. Różnica między dwiema a czterema warstwami przy typowej płytce prototypowej jest dziś rzędu kilkunastu dolarów przy pięciu sztukach, a przy większych nakładach maleje względnie. Konkretnie liczby dla trzech przypadków referencyjnych – płytki jednowarstwowej, dwuwarstwowej i czterowarstwowej – podamy w części trzeciej, przy omawianiu zamawiania produkcji PCB.

Prawdziwym kosztem czwartej warstwy nie są więc pieniądze, lecz utrata wymówki. Konstruktor, który dotąd

tłumaczył szумы w układzie ograniczeniami dwuwarstwowej płytki, po przejściu na cztery warstwy musi zmierzyć się z tym, że problem był gdzie indziej.

Co daje czwarta warstwa? Jedno, ale zasadnicze: ciągłą płaszczyznę odniesienia bezpośrednio pod każdą ścieżką sygnałową. Wszystko, co powiemy w części drugiej o drogach powrotnych prądu, o impedancji i o emisji, sprowadza się do tej jednej właściwości.

5.1. Jak wygląda realny stos czterowarstwowy

Popularny stos standardowy jednego z wytwórców azjatyckich, oznaczany symbolem JLC04161H-7628, dla płytki o grubości 1,6 mm ma następującą budowę: miedź warstwy zewnętrznej 0,035 mm, preimpregnat 7628 o grubości 0,2104 mm, miedź warstwy wewnętrznej 0,0152 mm, rdzeń 1,065 mm, miedź warstwy wewnętrznej 0,0152 mm, preimpregnat 7628 o grubości 0,2104 mm, miedź warstwy zewnętrznej 0,035 mm.

Z tych liczb wynikają dwa wnioski konstrukcyjne, które warto sobie przyswoić, zanim zacznie się cokolwiek trasować.

Po pierwsze, odległość od ścieżki na warstwie zewnętrznej do najbliższej płaszczyzny wynosi około 0,21 mm. To jest bardzo dobre sprzężenie: droga powrotna prądu jest krótka, pętla mała, emisja niska. Warstwy zewnętrzne są więc naturalnym miejscem dla sygnałów, pod warunkiem że warstwy drugą i trzecią zajmują ciągłą płaszczyznę.

Po drugie, odległość między dwiema płaszczyznami wewnętrzными wynosi 1,065 mm, czyli pięciokrotnie więcej. Pojemność międzypłaszczyznowa takiego stosu jest znikoma. Kto liczy na to, że płytka sama w sobie będzie kondensatorem odsprężającym dla zasilania, przy tym stosie się rozczaruje. Odsprężenie trzeba zrobić elementami i zrobić je porządnie.

Wniosek praktyczny co do przydziału warstw: warstwa pierwsza sygnałowa, druga masa, trzecia zasilanie lub druga masa, czwarta sygnałowa. Przydział, w którym płaszczyzny są na zewnątrz, a sygnały w środku, ma swoje uzasadnienie w ekranowaniu, ale przy tym stosie pogarsza sprężenie i utrudnia serwis.

5.2. Materiał i to, co się o nim zwykle błędnie zakłada

Przenikalność względna laminatu nie jest jedną liczbą i nie wynosi 4,3, choć taka wartość bywa wpisywana odruchowo. U cytowanego wytwórcy podano: 4,5 dla płytki dwuwarstwowej, 4,4 dla preimpregnatu 7628, 4,1 dla preimpregnatu 3313 i 4,16 dla preimpregnatu 2116. Maski lutownicze mają przenikalność 3,8 przy grubości co najmniej 10 mikrometrów, co obniża impedancję ścieżki mikropaskowej o kilka procent. Do obliczeń impedancji trzeba brać wartość dla konkretnego preimpregnatu w konkretnym stosie, a nie wartość przybliżoną.

Tolerancje, które trzeba znać: grubość gotowej płytki mieści się w granicach dziesięciu procent wartości nominalnej dla płytek o grubości co najmniej 1,0 mm – dla płytki 1,6 mm daje to zakres od 1,44 do 1,76 mm. Tolerancja impedancji kontrolowanej wynosi standardowo dziesięć procent, a zacieśnienie do pięciu procent jest możliwe na życzenie i wymaga materiałów o kontrolowanej grubości dielektryka oraz dodatkowej weryfikacji.

Symetria stosu nie jest kwestią estetyki. Płytki o niesymetrycznym rozkładzie miedzi i preimpregnatów wypacza się po lutowaniu rozpliwowym, co przy większych formatach uniemożliwia automatyczny montaż. Jeżeli jedna strona jest gęsto pokryta miedzią, a druga niemal pusta, warto wypełnić pustą stronę miedzią połączoną z masą.

5.3. Kiedy dwie warstwy wystarczą

Płytki dwuwarstwowa pozostaje rozsądnym wyborem w układach zasilania małej mocy, w prostej automatyce, w sterownikach z jednym mikrokontrolerem bez szybkich interfejsów oraz wszędzie tam, gdzie liczy się koszt jednostkowy przy dużej serii. Warunek jest jeden: strona spódni ma być w możliwie największym stopniu ciągłą masą, a rozmieszczenie elementów powinno pozwolić na trasowanie niemal wyłącznie po stronie górnej. Każda ścieżka poprowadzona po spodzie

Podłoże aluminiowe

Płytki jednowarstwowa na podłożu aluminiowym to osobna klasa wyrobów, spotykana przede wszystkim w oświetleniu i w energoelektronice małej mocy. Konstrukcja to warstwa miedzi, cienka warstwa dielektryka o podwyższonej przewodności cieplnej i podłoże aluminiowe pełniące funkcję radiatora. Konsekwencje projektowe: brak otworów metalizowanych, minimalna średnica otworu znacznie większa niż w laminacie zwykłym, wykończenie powierzchni ograniczone zwykle do cyny natryskowej, brak możliwości wykonania połączeń między stronami. Rozwiązanie ma sens, gdy trzeba odprowadzić ciepło z diod mocy albo tranzystorów bez osobnego radiatora, a schemat da się zrealizować na jednej warstwie miedzi z ewentualnymi zworami. Dobór dielektryka o odpowiedniej przewodności cieplnej omówimy w części trzeciej, przy materiałach.

przecina masę i tworzy szczelinę, którą prąd powrotny musi obejść.

6. Skąd biorą się liczby w regułach projektowych

Zestaw reguł projektowych nie jest zbiorem prawd uniwersalnych. Składa się z dwóch niezależnych warstw: z przepisanej tabeli możliwości technologicznych konkretnego wytwórcy oraz z wymagań elektrycznych wynikających z norm i z układu. Mylenie tych warstw jest przyczyną obu typowych błędów – projektowania na granicy możliwości wytwórcy tam, gdzie nie było takiej potrzeby, oraz przyjmowania domyślnych wartości narzędzia tam, gdzie potrzebna była decyzja.

6.1. Warstwa pierwsza: możliwości wytwórcy

Poniżej wybrane wartości z tabeli możliwości jednego z wytwórców azjatyckich, dla płytek FR-4 sztywnych. Traktujemy je jako przykład metody, nie jako wartości uniwersalne – u każdego wytwórcy tabela wygląda inaczej i trzeba ją przepisać osobno.

Z tabeli 2 warto wyprowadzić trzy praktyczne wnioski.

Tolerancja szerokości ścieżki wynosi dwadzieścia procent, więc zaprojektowana ścieżka o szerokości 0,1 mm może w rzeczywistości mieć od 0,08 do 0,12 mm. Dla ścieżki sygnałowej nie ma to znaczenia. Dla ścieżki zasilającej dobranej na styk do obliczonego prądu ma znaczenie zasadnicze, a dla ścieżki o kontrolowanej impedancji oznacza, że kontrolę impedancji trzeba zamówić osobno – wytwórca skoryguje wtedy szerokość w procesie, żeby trafić w zadaną wartość.

Tolerancja położenia otworu wynosi pięć setnych milimetra, a tolerancja samej średnicy trzynaście setnych w plus. Stąd biorą się minimalne pierścienie: to jest zapas na to, żeby przy najgorszym możliwym zbiegu odchylek otwór nie wyszedł poza pole. Projektowanie pierścienia o wartości bezwzględnego minimum oznacza zgodę na pewien odsetek płytek wadliwych.

Tabela 2.

Parametr	Wartość
Minimalna szerokość i odstęp ścieżki, miedź 1 oz	0,10/0,10 mm dla płytek 1- i 2-warstwowych; 0,09/0,09 mm dla wielowarstwowych
Minimalna szerokość i odstęp ścieżki, miedź 2 oz	0,16/0,16 mm dla płytek 2-warstwowych; 0,15/0,15 mm dla wielowarstwowych
Tolerancja szerokości ścieżki	±20 procent
Minimalna średnica wiercenia	0,15 mm dla płytek 2- i więcej warstwowych; 0,30 mm dla jednowarstwowych
Tolerancja otworu metalizowanego	+0,13/-0,08 mm
Tolerancja położenia otworu	±0,05 mm
Minimalny pierścień otworu metalizowanego	zalecane 0,20 mm, bezwzględne minimum 0,15 mm dla płytek wielowarstwowych, miedź 1 oz
Minimalny pierścień otworu niemetalizowanego	0,45 mm
Odstęp przelotki od ścieżki	0,20 mm
Odstęp otworu metalizowanego od ścieżki	zalecane 0,35 mm, minimum 0,28 mm
Odstęp miedzi od frezowanej krawędzi płytki	0,20 mm
Odstęp miedzi od krawędzi nacinanej metodą V-cut	0,40 mm
Przelotki ślepe i zagrzebane	nieobstugiwane, wyłącznie otwory na wylot

Nie należy projektować na granicy możliwości. Marża dwudziestu do trzydziestu procent nad wartościami minimalnymi kosztuje zwykle niewiele albo nic, a przekłada się bezpośrednio na uzysk. Zejście do wartości granicznych bywa dodatkowo płatne – u cytowanego wytwórcy otwory o średnicy 0,15 mm oraz otwory 0,20 i 0,25 mm z małym pierścieniem podlegają dopłacie.

6.2. Warstwa druga: wymagania elektryczne

Podstawowym dokumentem jest IPC-2221C, wydany w grudniu 2023 roku i zastępujący w całości rewizję B. Rewizja C nie jest kosmetyczna. Wprowadza między innymi nowe wytyczne co do doboru materiałów i folii miedzianej, paletyzację płytek, metalizację krawędzi,

przeregowane wymagania na minimalne odstępy izolacyjne z uwzględnieniem wytrzymałości dielektrycznej, efektu koronowego i upływności powierzchniowej, tabelę tolerancji impedancji, obszary wolne w płaszczyznach, tolerancje położenia cech, wyprowadzenia wciskane oraz nawiercanie wtórne.

Dla konstruktora najważniejsza jest zmiana dotycząca odstępów izolacyjnych. Wymagana odległość między przewodnikami zależy nie tylko od napięcia, lecz także od tego, czy przewodniki są na warstwie zewnętrznej czy wewnętrznej, czy są pokryte lakierem, oraz od wysokości nad poziomem morza – co ma praktyczne znaczenie w urządzeniach transportowanych lotniczo i pracujących w górach. Wartości tabelaryczne podaje norma i nie wolno ich tutaj przytaczać; dostęp do dokumentu jest odpłatny.

Tabela możliwości wytwórcy

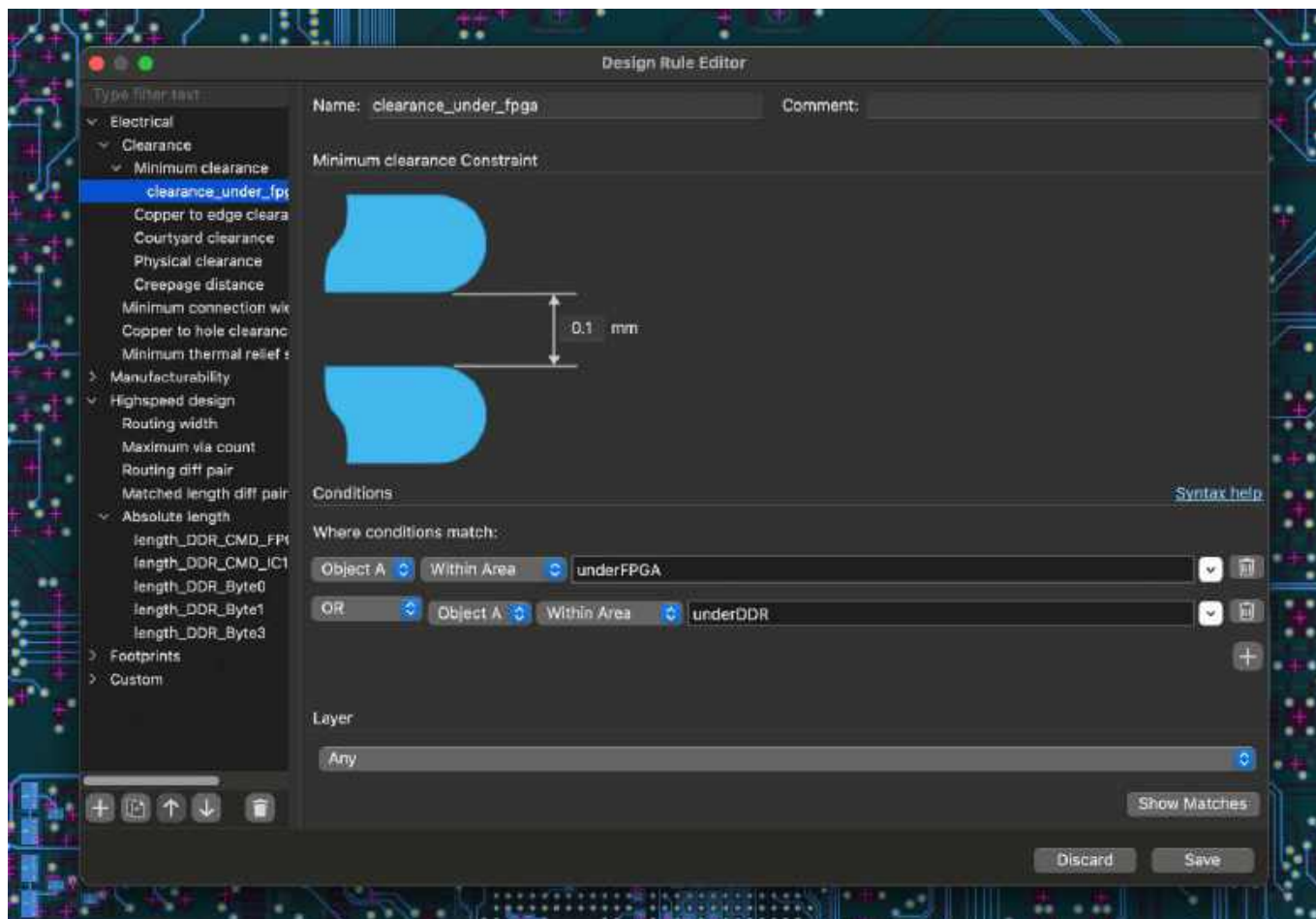
Reguły projektowe w narzędziu

(wartość z marginesem)

Ścieżka / odstęp, 1 oz: 0,10 / 0,10 mm	→	Clearance / Track width 0,13 mm
Minimalna średnica wiercenia: 0,15 mm	→	Hole size (min) 0,20 mm
Pierścień otworu metalizowanego: 0,20 mm	→	Annular ring (min) 0,25 mm
Odstęp przelotki od ścieżki: 0,20 mm	→	Via to track clearance 0,25 mm
Miedź – frezowana krawędź: 0,20 mm	→	Copper to board edge 0,25 mm
Mostek maski (zielona): 0,10 mm	→	Solder mask sliver / bridge 0,13 mm

margines 20–30 % nad wartością minimalną

Rysunek 4. Przepisanie tabeli możliwości wytwórcy na reguły projektowe w narzędziu, z marginesem nad wartością minimalną. Rysunek własny; wartości za tabelą możliwości – jlcpcb.com/capabilities/pcb-capabilities



Rysunek 5. Graficzny edytor reguł projektowych wprowadzony w KiCadzie 10. Ilustracja źródłowa: kicad.org/img/blog/2026/version-10-release/design_rule_editor.png (KiCad, licencja CC BY 3.0)

6.3. Obciążalność prądowa: dlaczego popularne kalkulatory zawyżają szerokość

Większość dostępnych w internecie kalkulatorów szerokości ścieżki realizuje wzór z IPC-2221, w postaci: prąd równa się współczynnik k razy przyrost temperatury do potęgi 0,44 razy pole przekroju do potęgi 0,725, przy polu przekroju wyrażonym w tysięcznych cali kwadratowych. Współczynnik k wynosi 0,048 dla ścieżek na warstwach zewnętrznych i 0,024 dla ścieżek na warstwach wewnętrznych – warstwa wewnętrzna, otoczona laminatem o niskiej przewodności cieplnej, oddaje ciepło mniej więcej dwukrotnie gorzej niż ścieżka chłodzona powietrzem. Kalkulator z tym wzorem udostępniają między innymi Digi-Key ([digikey.pl/en/resources/conversion-calculators/conversion-calculator-pcb-trace-width](https://www.digikey.com/en/resources/conversion-calculators/conversion-calculator-pcb-trace-width)) oraz LCSC ([lcsc.com/tools/conversion-calculator-pcb-trace-width](https://www.lcsc.com/tools/conversion-calculator-pcb-trace-width)).

Dla miedzi 1 oz, czyli o grubości około 35 mikrometrów, ścieżki o szerokości 0,254 mm i dla przyrostu temperatury 10 K wzór daje w przybliżeniu 0,9 A dla ścieżki na warstwie zewnętrznej.

Problem z tym wzorem jest natury źródłowej. Dane, na których go oparto, pochodzą z badań z lat pięćdziesiątych dwudziestego wieku, wykonanych na płytkach

jedno- i dwuwarstwowych z materiałów, których się dziś nie stosuje. Wzór nie uwzględnia obecności płaszczyzn miedzi, grubości płytki, przewodności cieplnej laminatu ani sąsiedztwa dużych obszarów miedzi. W efekcie jest zachowawczy: prowadzi do ścieżek szerszych, niż potrzeba.

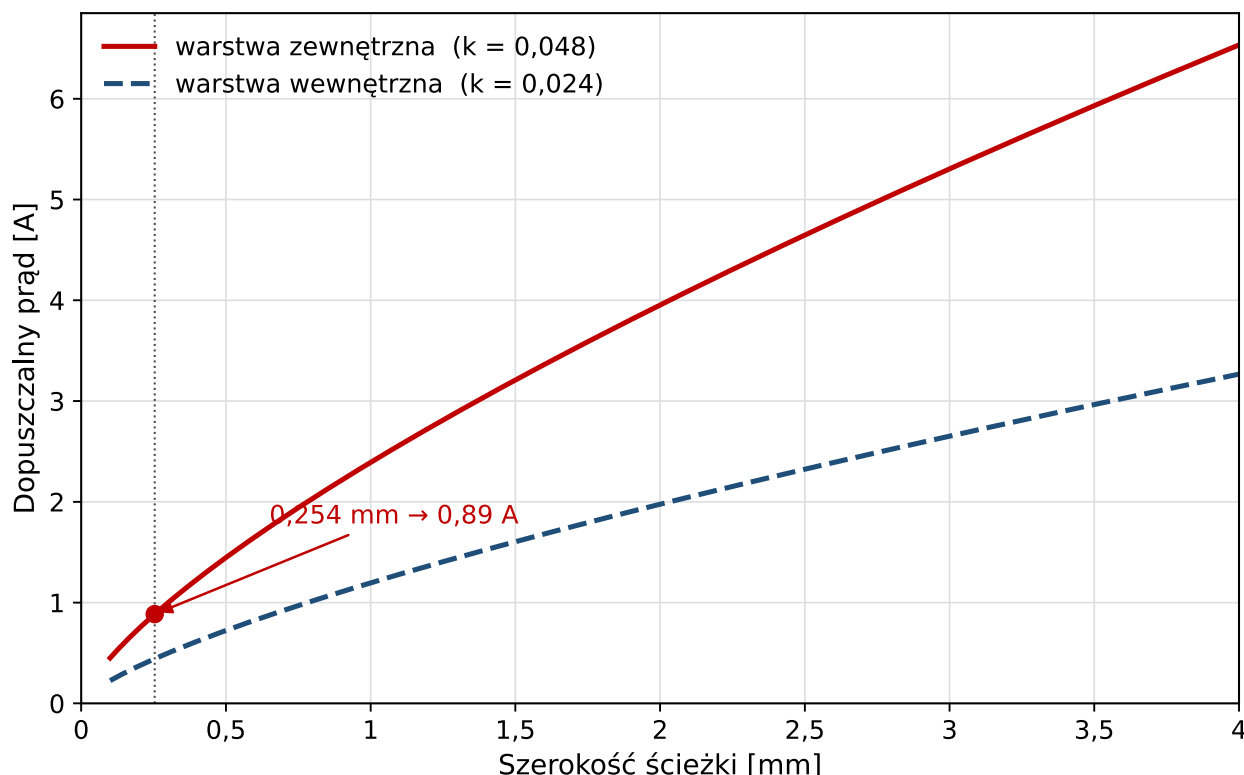
Dokumentem właściwym jest IPC-2152, wydany w 2009 roku i oparty na nowych badaniach, uwzględniający wpływ płaszczyzn, materiału i geometrii. Sama norma IPC-2221C w rewizji C odsyła do IPC-2152 jako do źródła preferowanego dla obliczeń obciążalności prądowej, zachowując własne wykresy jedynie jako szybkie sprawdzenie zachowawcze.

Zalecenie praktyczne. Do ścieżek sygnałowych wzór z IPC-2221 wystarczy z nadmiarem. Do ścieżek mocy, zwłaszcza w przetwornicach i w torach zasilania powyżej kilku amperów, należy korzystać z IPC-2152 albo z narzędzia, które ten dokument implementuje. Różnica potrafi wynosić kilkadziesiąt procent szerokości, a w gęstym projekcie to jest różnica między płytką czterowarstwową a sześciowarstwową.

7. Rozmieszczenie i trasowanie

Osiemdziesiąt procent jakości płytki powstaje na etapie rozmieszczenia elementów. Trasowanie tej jakości nie

Obciążalność prądowa według wzoru IPC-2221 miedź 1 oz (35 μm), przyrost temperatury 10 K



Rysunek 6. Obciążalność prądowa według wzoru z IPC-2221 dla miedzi 1 oz przy przyroście temperatury 10 K, dla warstwy zewnętrznej ($k = 0,048$) i wewnętrznej ($k = 0,024$). Rysunek własny z wzoru. Danych pomiarowych IPC-2152 nie odwzorowano – dokument jest chroniony prawem autorskim; różnicę omówiono w tekście

tworzy – ono ją tylko ujawnia. Projekt, w którym trasowanie jest udręką, jest zwykle projektem źle rozmieszczonym, a nie projektem trudnym.

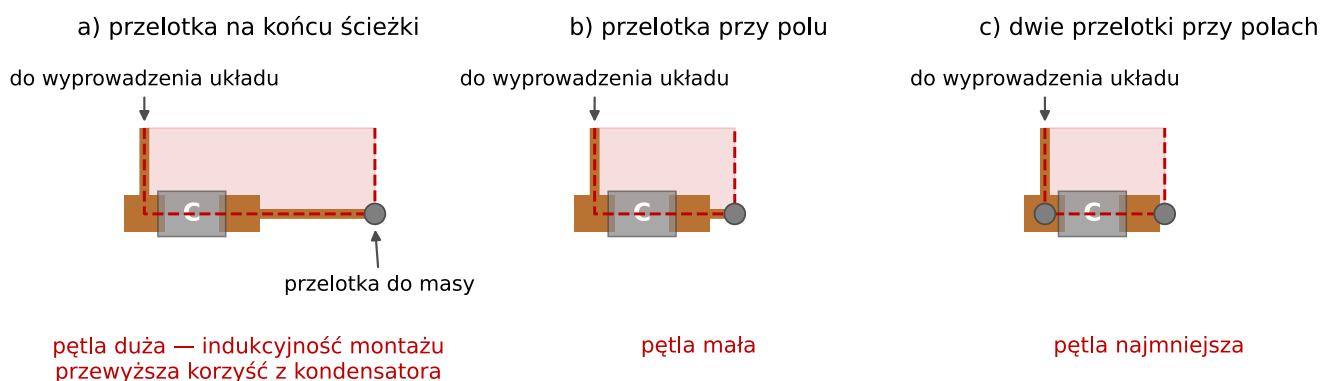
7.1. Kolejność działań

Kolejność, od której nie warto odstępować:

- Ograniczenia mechaniczne: obrys płytki, otwory montażowe, strefy zakazane, wysokości elementów. Import geometrii z modelu obudowy, nie przepisywanie wymiarów z rysunku.
- Złącza i elementy o narzuconym położeniu: przyciski, diody sygnalizacyjne, gniazda, elementy dostępne od zewnątrz.

- Bloki zasilania, w kolejności przepływu energii: wejście, zabezpieczenia, filtr, przetwornica, odbiorniki.
- Układy szybkie i wrażliwe, wraz z ich najbliższym otoczeniem: rezonatory, przetworniki analogowo-cyfrowe, interfejsy różnicowe.
- Cała reszta.

Rozmieszczenie powinno odwzorowywać schemat blokowy urządzenia, nie estetykę obwodów. Elementy jednego bloku funkcjonalnego trzymamy razem, bloki rozdzielamy, przepływ sygnału prowadzimy w jednym kierunku. Płytkę, która wygląda jak uporządkowany schemat blokowy, trasuje się sama.



Rysunek 7. Pętla prądowa kondensatora odsprężającego w trzech wariantach montażu przelotki masowej. O skuteczności decyduje pole pętli, nie odległość kondensatora od wyprowadzenia układu

7.2. Zasilanie i odsprężanie

Droga od złącza zasilania przez zabezpieczenia, filtr i przetwornicę do odbiornika ma być krótka i jednokierunkowa. Zawrócenie tej drogi oznacza, że prąd wejściowy i wyjściowy przetwornicy płyną obok siebie, a filtr wejściowy przestaje działać w sposób, który zaprojektowano.

Przy kondensatorach odsprężających liczy się pole pętli, nie sama odległość od wyprowadzenia. Kondensator umieszczony blisko wyprowadzenia, ale połączony z masą przez ścieżkę o długości pięciu milimetrów i przelotkę oddaloną o kolejne trzy, ma indukcyjność montażu przewyższającą to, co daje jego pojemność. Przelotka do masy należy do pola lutowniczego kondensatora, a nie do ścieżki wychodzącej z tego pola. Przy obudowach o rastrze pozwalającym na to warto stosować dwie przelotki.

Dobór wartości: kilka kondensatorów o jednakowej wartości rozmieszczonych blisko wyprowadzeń zasilania działa lepiej niż drabinka wartości od jednego nanofarada do dziesięciu mikrofaradów, rozłożona po całej płytce. Do zagadnienia impedancji docelowej sieci dystrybucji zasilania i do zjawiska antyrezonansu między kondensatorami o różnych wartościach wrócimy w części drugiej.

7.3. Masa

Jedna ciągła płaszczyzna masy. To jest reguła, od której odstępstwa wymagają uzasadnienia, a nie odwrotnie.

Praktyka dzielenia masy na „analogową” i „cyfrową”, z połączeniem ich w jednym punkcie, wyrządziła w ciągu ostatnich trzydziestu lat więcej szkody niż pożytku i jest

dziś powszechnie krytykowana w literaturze przedmiotu. Powód jest prosty: przy częstotliwościach powyżej kilkadziesiąt kiloherców prąd powrotny nie płynie najkrótszą drogą, tylko drogą o najmniejszej indukcyjności, czyli bezpośrednio pod ścieżką sygnałową. Rozcięcie płaszczyzny zmusza ten prąd do objazdu, a objazd jest pętlą, czyli anteną. Zagadnienie jest na tyle ważne, że poświęcimy mu znaczną część odcinka drugiego.

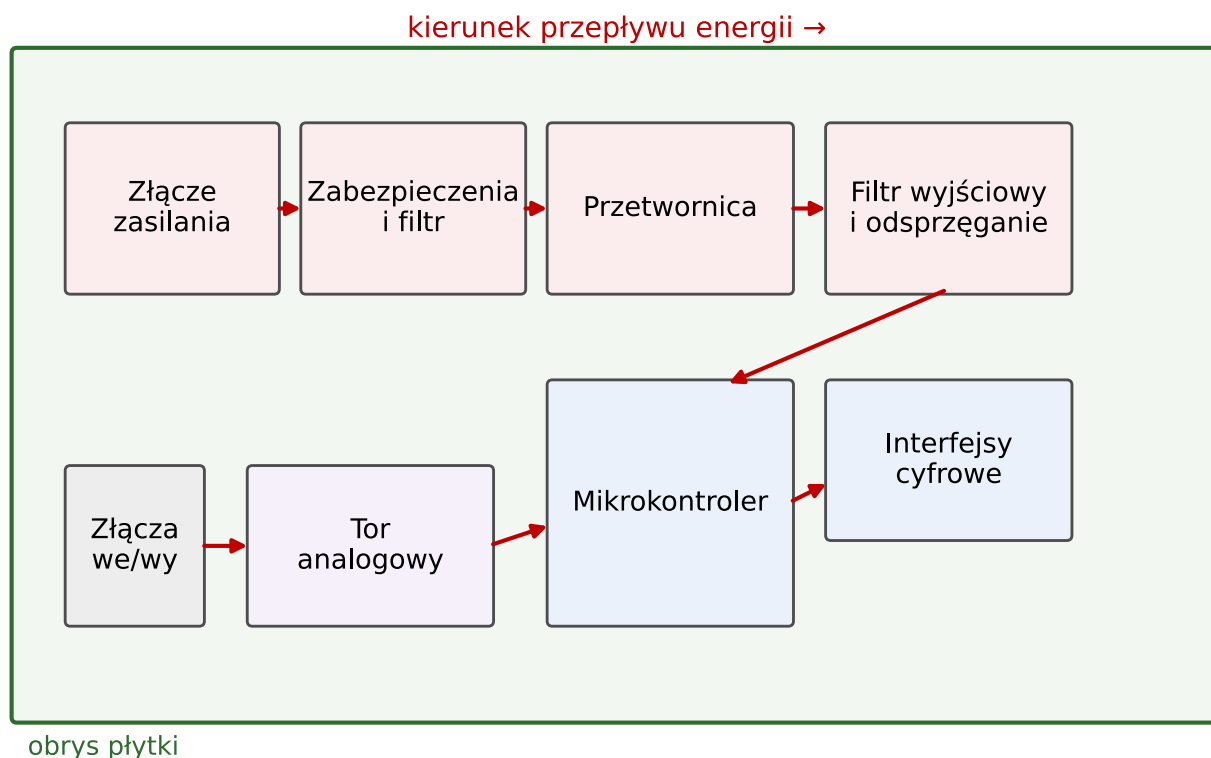
Wniosek na teraz: nie rozcinać masy, dopóki nie potrafisz narysować, którędy popłynie prąd powrotny każdego sygnału przecinającego rozcięcie.

Zmiana warstwy przez przelotkę oznacza zmianę płaszczyzny odniesienia, a więc przerwę w drodze powrotnej. Przy sygnałach szybkich obok przelotki sygnałowej powinna znaleźć się przelotka łącząca obie płaszczyzny odniesienia. Przy stosie omawianym w punkcie piątym, gdzie warstwy druga i trzecia są oddalone o milimetr, ma to znaczenie większe, niż mogłoby się wydawać.

7.4. Trasowanie

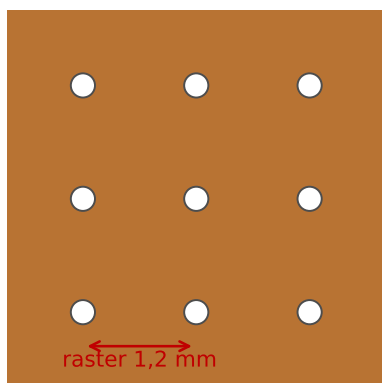
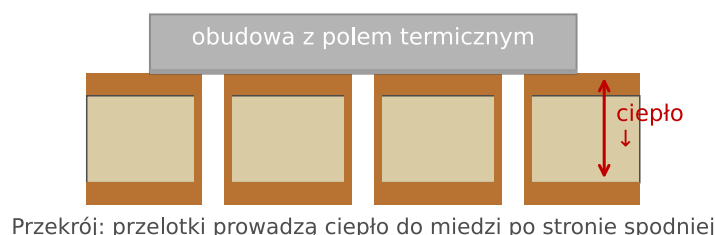
Najpierw sygnały krytyczne – zegary, interfejsy różnicowe, wejścia analogowe o wysokiej impedancji, tory mocy. Potem reszta. Autorouter nadaje się do połączeń nieistotnych i do sprawdzenia, czy rozmieszczenie w ogóle da się rozprowadzić. Nie nadaje się do sygnałów, których przebieg trzeba kontrolować.

Szerokości domyślne: dla sygnałów logicznych 0,15 do 0,20 mm jest wartością rozsądną, wyraźnie powyżej minimum technologicznego, a więc bezpieczną dla



Rysunek 8. Rozmieszczenie odwzorowujące schemat blokowy urządzenia: tor zasilania prowadzony jednokierunkowo, bloki funkcjonalne rozdzielone, tor analogowy odsunięty od przetwornicy

Pole termiczne od strony montażu

przelotki $\varnothing$ 0,3 mm, matryca 3 × 3

Przekrój: przelotki prowadzą ciepło do miedzi po stronie spodniej

Rysunek 9. Przelotki termiczne pod obudową z polem masowym – widok od strony montażu i przekrój

uzysku. Dla ścieżek zasilających szerokość wynika z obliczenia, o którym była mowa w punkcie szóstym, nigdy z oszacowania na oko.

O kątach prostych warto powiedzieć wprost, bo mit trwa. Przy częstotliwościach, z jakimi ma do czynienia większość konstruktorów, kąt prosty na ścieżce nie ma mierzalnego wpływu na sygnał. Powód, dla którego się go unika, jest technologiczny: ostry kąt wewnętrzny bywa pułapką dla wytrawiacza i miejscem lokalnego podtrawienia. To wystarczający powód, żeby stosować kąty czterdziestopięciostopniowe, ale nie jest to powód, dla którego zwykle się je zaleca.

Wylewane obszary miedzi wymagają kontroli. Wysepki miedzi bez połączenia z żadnym obwodem należy usunąć – są anteną i nic poza tym. Wąskie przewężenia w wylewce, przez które ma płynąć prąd powrotny, trzeba obejrzeć osobno, ponieważ kontrola DRC ich nie zgłosi.

8. Termika i ścieżki mocy

Ciepło jest w płytce drukowanej odprowadzane przede wszystkim przez miedź, dopiero potem przez laminat i konwekcję. Z tego wynika kilka reguł praktycznych.

Pod obudowami z polem termicznym na spodzie – a więc pod większością współczesnych obudów mocy i pod znaczną częścią układów cyfrowych – konieczne są przelotki termiczne. Ich liczba i średnica są kompromisem między odprowadzaniem ciepła a ryzykiem zassania pasty lutowniczej do otworu podczas rozplwy, co skutkuje pustką w spoinie i pogorszeniem zarówno przewodzenia ciepła, jak i mocowania.

Rozwiązania są dwa. Przy przelotkach zamykanych maską lutowniczą u cytowanego wytwórcy obowiązują ograniczenia: przelotka wypełniana maską nie może mieć odsłony maski po żadnej ze stron, jej średnica nie może przekraczać 0,5 mm, a odstęp od innych odsłon maski, w tym od pól lutowniczych, musi wynosić co najmniej 0,35 mm. Przy gęstych obudowach z polem

termicznym te warunki bywają trudne do spełnienia jednocześnie.

Rozwiązaniem drugim jest przelotka w polu lutowniczym, wypełniona żywicą epoksydową albo pastą miedzianą i pokryta warstwą miedzi. U cytowanego wytwórcy proces ten jest domyślny dla płytek sześciowarstwowych i wyższych i obsługuje średnice przelotek od 0,15 do 0,55 mm, przy czym wypełnienie pastą miedzianą jest przeznaczone do zastosowań o podwyższonych wymaganiach cieplnych. Na płycie czterowarstwowej jest to opcja dodatkowo płatna, którą trzeba zamówić świadomie.

Miedź jako radiator: przejście z 1 oz na 2 oz podwaja przekrój i wyraźnie poprawia rozprowadzanie ciepła w płaszczyźnie, ale kosztuje. Kosztuje nie tylko w cenie płytki – przy miedzi 2 oz minimalna szerokość i odstęp ścieżki rosną z 0,10 do 0,16 mm dla płytek dwuwarstwowych, a minimalny mostek maski między polami do 0,20 mm niezależnie od koloru. Decyzja o grubszej miedzi jest więc decyzją o zagęszczeniu całego projektu i trzeba ją podjąć przed rozmieszczeniem, a nie po nim.

Elementy ciepłe trzymamy z dala od elementów wrażliwych termicznie: od rezonatorów kwarcowych, od precyzyjnych źródeł napięcia odniesienia, od czujników temperatury i od kondensatorów elektrolitycznych, których trwałość spada w tempie wykładniczym wraz z temperaturą pracy. Jeżeli rozmieszczenie nie pozwala na odsunięcie, wycięcie w laminacie albo szczelina bez miedzi między obszarami ogranicza przepływ ciepła w płaszczyźnie.

Właściwości cieplne samego laminatu – temperatura zeszklenia, temperatura rozkładu, przewodność cieplna, a także odporność na tworzenie włókien anodowych – omówimy w części trzeciej, przy doborze materiału.

9. Dane wyjściowe: co naprawdę wysyłamy do wytwórcy

Projekt kończy się nie w chwili domknięcia ostatniej ścieżki, lecz w chwili, gdy pakiet danych produkcyjnych

zostaje sprawdzony w narzędziu niezależnym od tego, w którym powstał.

9.1. Format

Obowiązującym formatem jest Gerber w wersji rozszerzonej, w rewizji 2024.05 opublikowanej 3 lipca 2024 roku przez firmę Ucamco. Standard Gerber, czyli RS-274-D, jest przestarzały i formalnie wycofany – nie należy go używać, mimo że część narzędzi wciąż potrafi go wyeksportować.

Rozróżnienie, które warto rozumieć: pliki z atrybutami to Gerber X2, pliki bez atrybutów to X1. Atrybuty są etykietami niosącymi informacje o znaczeniu obiektów, bez wpływu na sam obraz: która warstwa jest górną maską, a która dolną miedzią, czy dany błysk jest polem SMD, przelotką czy znacznikiem bazowym, do której sieci należy dane połączenie. Wytwórca, który dostaje X2, nie musi zgadywać przyporządkowania warstw, a jego kontrola przedprodukcyjna wyłapie więcej. Eksport z atrybutami jest w większości narzędzi jedną opcją do zaznaczenia i nie ma powodu, żeby jej nie zaznaczyć.

Alternatywami jednoplikowymi są IPC-2581 i ODB++. Oba niosą w jednym pliku pełen opis produkcji łącznie ze stosem warstw, netlistą i wymaganiami. Oba są przez większych wytwórców obsługiwane. Praktyka pozostaje jednak taka, że pakiet gerberowy przyjmie każdy, a plik IPC-2581 nie każdy – dlatego zalecenie brzmi: wysyłaj Gerbery X2, a plik IPC-2581 albo ODB++ dołącz jako materiał uzupełniający.

9.2. Zawartość pakietu

Kompletny pakiet produkcyjny gołej płytki zawiera:

- warstwy miedzi, wszystkie, w oznaczeniu jednoznacznie wskazującym kolejność w stosie;

Pakiet gerberowy (Gerber X2)

kilkaście plików, atrybuty niosą znaczenie warstw

Warstwy miedzi (L1...Ln)
Maska górna i dolna
Opis montażowy
Obrys i frezowanie
Wiercenia PTH
Wiercenia NPTH
Gerber Job (.gbrjob)
Netlista IPC-D-356
Rysunek technologiczny

Przyjmuje każdy wytwórca

ta sama
treść

IPC-2581 albo ODB++

jeden plik, ta sama zawartość plus stos warstw

- Obrazy wszystkich warstw
- Wiercenia
- Stos warstw i materiały
- Netlista
- Wymagania impedancyjne
- Lista materiałowa
- Rozmieszczenie elementów

Przyjmują wytwórcy więksi

- maskę lutowniczą górną i dolną;
- opis montażowy górny i dolny;
- obrys płytki wraz z otworami frezowanymi i szczelinami;
- pliki wierceń, osobno dla otworów metalizowanych i niemetalizowanych;
- plik zadania Gerber Job albo równoważny dokument opisujący stos, grubość, wykończenie, kolor maski i wymagania dotyczące impedancji;
- rysunek technologiczny;
- netlistę w formacie IPC-D-356, umożliwiającą wytwórcy elektryczną kontrolę zgodności z projektem.

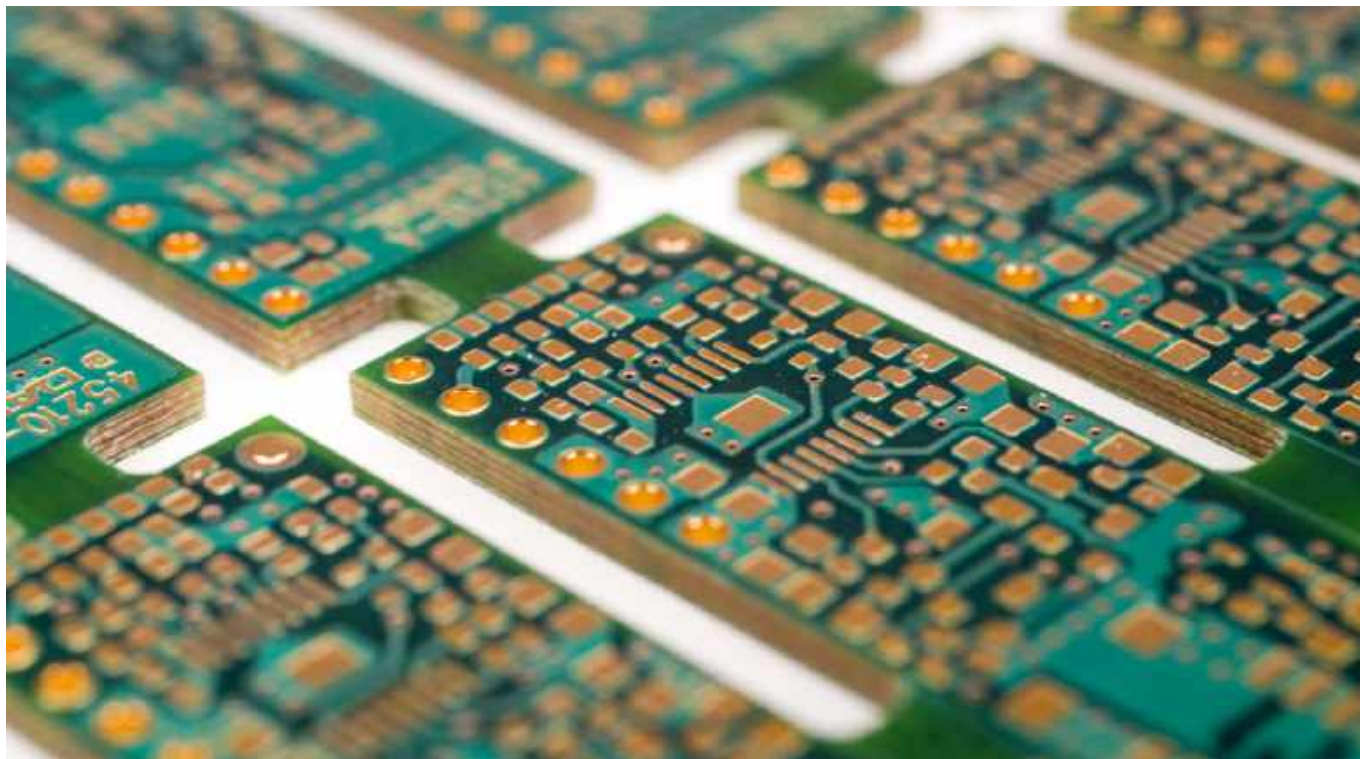
Do montażu dochodzą lista materiałowa i plik rozmieszczenia elementów, którym poświęcimy osobne miejsce w części czwartej.

Rysunek technologiczny bywa pomijany i jest to błąd. Powinien zawierać wymiary gabarytowe z tolerancjami, przekrój stosu z materiałami i grubościami, tabelę wierceń, wymagania co do wykończenia powierzchni, wymagania co do impedancji z podaniem warstw i szerokości, klasę wykonania według IPC, wymagania dotyczące znakowania oraz informację o panelizacji. Wszystko, czego nie ma na rysunku, wytwórca rozstrzygnie po swoim.

9.3. Weryfikacja

Zasada jest jedna i nie ma od niej wyjątku: pakiet trzeba obejrzeć w przeglądarce niezależnej od narzędzia projektowego. Podgląd wbudowany pokazuje to, co narzędzie zamierzało wyeksportować. Przeglądarka niezależna pokazuje to, co rzeczywiście znalazło się w plikach – a to są dwie różne rzeczy, o czym przekonuje się każdy, kto choć raz wysłał pakiet z pominiętą warstwą wierceń niemetalizowanych.

Rysunek 10. Zawartość pakietu produkcyjnego: pakiet gerberowy wobec pliku IPC-2581. Rysunek własny; struktura za specyfikacją formatu Gerber, rewizja 2024.05 – ucamco.com/en/gerber/downloads



Bezpłatną przeglądarkę referencyjną udostępnia Ucamco, autor formatu; jej zaletą jest ścisła zgodność ze specyfikacją i formalna kontrola składni, dzięki czemu wychwytuje ona konstrukcje, które inne przeglądarki interpretują pobłażliwie, a wytwórca niekoniecznie. Wytwórcy azjatyccy udostępniają dodatkowo własne narzędzia kontroli wykonalności, działające w przeglądarce internetowej i zwracające listę problemów w odniesieniu do konkretnej tabeli możliwości. Warto skorzystać z obu rodzajów narzędzi – sprawdzają co innego.

Archiwizacja: pakiet produkcyjny każdego wydania płytki przechowujemy razem z sumą kontrolną i datą, w postaci nienaruszonej. Kiedy po dwóch latach trzeba zrobić dwadzieścia sztuk, wysyła się archiwum, a nie eksportuje ponownie z narzędzia, które w międzyczasie zmieniło wersję i domyślne ustawienia.

10. Lista kontrolna przed wystaniem

Do wydrukowania i powieszenia nad biurkiem.

- ☐ Kontrola ERC przechodzi bez ostrzeżeń albo każde ostrzeżenie ma pisemne uzasadnienie.
- ☐ Każdy element ma w polach producenta, numer katalogowy, obudowę i parametr krytyczny.
- ☐ Każde pole lutownicze pochodzi z noty katalogowej producenta elementu albo z generatora zgodnego z IPC – żadne nie pochodzi z nieznanego źródła.
- ☐ Wydruk w skali jeden do jednego przyłożony do rzeczywistych elementów przynajmniej dla obudów nietypowych.
- ☐ Każdy element ma model trójwymiarowy; sprawdzono kolizje z obudową urządzenia.
- ☐ Oznaczenie wyprowadzenia pierwszego widoczne na opisie montażowym poza obrysem obudowy.
- ☐ Opis montażowy nie nachodzi na żadne pole lutownicze; wysokość znaków co najmniej 1,0 mm.
- ☐ Reguły projektowe przepisane z tabeli konkretnego wytwórcy, z marżą co najmniej dwudziestu procent nad wartościami minimalnymi.
- ☐ Sprawdzono minimalny mostek maski wobec wybranego koloru maski.
- ☐ Stos warstw zdefiniowany w narzędziu zgodnie z rzeczywistym stosem wytwórcy, nie domyślny.
- ☐ Płaszczyzny odniesienia ciągłe pod wszystkimi sygnałami szybkimi; obejrzano wylewki warstwa po warstwie.
- ☐ Przy każdej przelotce sygnału szybkiego zmieniającego warstwę odniesienia jest przelotka łącząca płaszczyzny.
- ☐ Każdy kondensator odsprężający ma własną przelotkę do masy przy polu lutowniczym.
- ☐ Szerokości ścieżek zasilających wynikają z obliczenia, z zapisanym przyrostem temperatury.
- ☐ Przelotki termiczne pod polami masowymi rozwiązane świadomie: maska albo wypełnienie z pokryciem.
- ☐ Usunięto wysepki miedzi bez połączenia.
- ☐ Zaplanowano punkty pomiarowe do uruchamiania.
- ☐ Kontrola DRC przechodzi bez błędów.
- ☐ Pakiet obejrzany w przeglądarce niezależnej od narzędzia projektowego oraz w narzędziu kontroli wykonalności wytwórcy.
- ☐ Rysunek technologiczny kompletny.

Słownik

Terminy uporządkowane alfabetycznie. Jednostce „1 oz”, która w tekście pojawia się najczęściej i budzi najwięcej nieporozumień, poświęcona jest osobna ramka na końcu.

Atrybuty (Gerber X2). Etykiety tekstowe dołączone do plików Gerber, niosące informację o znaczeniu obiektów bez wpływu na sam obraz: która warstwa jest górną maską, czy dany kształt jest polem lutowniczym, przelotką czy znacznikiem bazowym, do której sieci należy połączenie. Pliki bez atrybutów to Gerber X1, z atrybutami – X2.

Autorouter. Moduł narzędzia projektowego prowadzący połączenia automatycznie, na podstawie listy sieci i reguł projektowych.

BGA. Obudowa z wyprowadzeniami w postaci matrycy kulek lutowniczych pod spodem korpusu, niedostępnych do oględzin po montażu.

CAF. Zjawisko powstawania przewodzących włókien wzdłuż osnowy szklanej laminatu pod wpływem napięcia i wilgoci, prowadzące do upływności między sąsiednimi otworami metalizowanymi.

Droga powrotna. Trasa, którą prąd wraca od odbiornika do źródła. Przy częstotliwościach powyżej kilkudziesięciu kiloherców biegnie nie najkrótszą drogą, lecz drogą o najmniejszej indukcyjności, czyli bezpośrednio pod ścieżką sygnałową.

Dk (przenikalność względna). Właściwość dielektryka określająca, jak bardzo spowalnia on falę elektromagnetyczną. Wraz z geometrią ścieżki wyznacza jej impedancję charakterystyczną.

DFM (kontrola przedprodukcyjna). Weryfikacja projektu przez inżynierów wytwórcy przed uruchomieniem produkcji, wykrywająca cechy niemożliwe do wykonania w danym procesie – zbyt wąskie ścieżki, zbyt małe pierścienie, kolizje maski.

DRC (kontrola reguł projektowych). Automatyczne sprawdzenie layoutu wobec zdefiniowanego zestawu reguł: odstępów, szerokości, średnic otworów, stref zakazanych.

ENIG. Wykończenie powierzchni: chemiczny nikiel pokryty cienką warstwą złota zanurzeniowego. Powierzchnia płaska, trwała w przechowywaniu, droższa od cyny natryskowej.

ERC (kontrola reguł elektrycznych). Sprawdzenie schematu pod kątem błędów połączeń: niepodłączonych wejść, konfliktu dwóch wyjść w jednej sieci, zasilania bez źródła.

Format Gerber. Standardowy tekstowy format plików zawierający wektorowy opis poszczególnych warstw płytki – gdzie ma być miedź, gdzie odstona maski, gdzie opis montażowy. Obowiązuje wersja rozszerzona; wersja RS-274-D jest wycofana.

FR-4. Najpopularniejszy laminat konstrukcyjny: tkanina szklana nasyczona żywicą epoksydową o obniżonej palności. Oznaczenie opisuje klasę materiału, nie konkretny wyrób – parametry różnią się między producentami.

HASL. Wykończenie powierzchni polegające na zanurzeniu płytki w stopie lutowniczym i zdmuchnięciu nadmiaru gorącym powietrzem. Tanie, ale daje powierzchnię nierówną.

IPC-2581, ODB++. Formaty opisu produkcji zapisujące w jednym pliku obrazy wszystkich warstw, wiercenia, stos warstw, netlistę i wymagania. Alternatywa dla wieloplikowego pakietu gerberowego.

Impedancja kontrolowana. Wymaganie, by ścieżka o zadanej geometrii miała określoną impedancję charakterystyczną. Wytwórca koryguje wtedy szerokość ścieżki w procesie i potwierdza wynik pomiarem na kuponie technologicznym.

Laminat. Materiał konstrukcyjny płytki: dielektryk wzmocniony tkaniną szklaną, pokryty folią miedzianą.

Maska lutowicza. Polimerowa warstwa ochronna, najczęściej zielona, zabezpieczająca miedź przed utlenianiem i zapobiegająca zwarciom podczas lutowania. Odstony w masce wyznaczają miejsca lutowane.

Menedżer ograniczeń. Podsystem narzędzia projektowego pozwalający przypisać wymagania – impedancję, dopasowanie długości, odstęp – do sieci i klas sieci, a nie do pojedynczych ścieżek.

Mil. Tysięczna część cala, czyli 0,0254 mm. Jednostka używana w normach IPC i w tabelach wytwórców.

Mostek maski. Pasek maski lutowicznej pozostawiony między sąsiednimi polami lutowniczymi. Poniżej pewnej szerokości nie utrzymuje się na płycie i odpada w procesie.

Netlista. Tekstowy spis wszystkich połączeń elektrycznych w projekcie. W wersji IPC-D-356 służy wytwórcy do elektrycznej kontroli zgodności gotowej płytki z projektem.

OSP. Wykończenie powierzchni w postaci cienkiej powłoki organicznej chroniącej miedź. Najtańsze, o ograniczonym czasie przechowywania.

Panelizacja. Złożenie wielu egzemplarzy płytki w jeden arkusz produkcyjny, z ramkami technologicznymi i liniami rozdzielania.

Pierścień (annular ring). Obrączka miedzi wokół otworu metalizowanego. Jej szerokość jest zapasem na odchyłki położenia i średnicy wiercenia.

Pole lutownicze (pad). Odstoniony fragment miedzi, do którego lutuje się wyprowadzenie elementu. Pole SMD dotyczy elementów montowanych powierzchniowo.

Poziom gęstości pola lutowniczego. Wariant geometrii pola zdefiniowany w normach IPC: od największej ilości materiału, dającej najdłuższe pola i największy zapas na spoinę, po najmniejszą, przeznaczoną do rozmieszczeń najgęstszych.

Preimpregnat (prepreg). Tkanina szklana nasyczona żywicą w stanie nieutwardzonym, spajająca rdzenie i folie miedziane podczas prasowania stosu.

Przelotka (via). Metalizowany otwór łączący ścieżki na różnych warstwach. Przelotka ślepa łączy warstwę zewnętrzną z wewnętrzną, zagrzebana wyłącznie warstwy wewnętrzne, przelotka na wylot przechodzi przez całą płytkę.

Przelotka zszywająca. Przelotka łącząca dwie płaszczyzny odniesienia, umieszczana obok przelotki sygnałowej, by prąd powrotny miał ciągłą drogę przy zmianie warstw.

PTH, NPTH. Otwór metalizowany i otwór niemetalizowany. Wiercenia obu rodzajów wysyła się do wytwórcy w osobnych plikach.

QFN. Obudowa bez wyprowadzeń bocznych, z polami lutowniczymi na spodzie krawędzi i zwykle z polem termicznym pośrodku.

Rdzeń (core). Utwardzony laminat pokryty miedzią z obu stron, stanowiący sztywną podstawę stosu warstw.

Stos warstw (stackup). Opis struktury płytki: kolejność warstw miedzi i dielektryka, rodzaje materiałów i ich grubości.

Tg, Td. Temperatura zeszklenia i temperatura rozkładu laminatu. Pierwsza wyznacza granicę, powyżej której materiał mięknie, druga – powyżej której zaczyna się rozkładać chemicznie.

Via-in-pad. Przelotka umieszczona w polu lutowniczym, wypełniona żywicą albo pastą miedzianą i pokryta miedzią, tak by pole pozostało płaskie.

Wylewka miedzi (pour). Obszar miedzi wypełniający wolną przestrzeń warstwy, zwykle połączony z masą.

V-cut. Nacięcie klinowe wykonane z obu stron arkusza, pozwalające rozdzielić płytki przez wytłamanie.

Znacznik bazowy (fiducial). Punkt odniesienia na płycie, zwykle miedziane kółko z odstoną w masce, służący kamerom automatu montażowego do precyzyjnego ustalenia położenia płytki.

Skąd się wzięła jednostka „1 oz”

Grubość folii miedzianej na płycie podaje się nie w mikrometrach, lecz w uncjach – i nie jest to pomyłka jednostki. Zapis „1 oz” oznacza jedną uncję miedzi rozwalcowaną na powierzchnię jednej stopy kwadratowej. Uncja to jednostka masy, stopa kwadratowa jednostka powierzchni, a ich iloraz daje jednoznacznie określoną grubość warstwy.

Po przeliczeniu: 1 oz odpowiada grubości około 35 mikrometrów, czyli 1,378 mila. Stąd wartości spotykane w tabelach wytwórców: 0,5 oz to około 17,5 mikrometra, 2 oz to około 70 mikrometrów, 3 oz około 105 mikrometrów. Zależność jest liniowa, więc przeliczenie sprowadza się do mnożenia.

Jednostka pochodzi z amerykańskiej praktyki wytwórczej i utrwaliła się dlatego, że producent laminatu sprzedaje folię miedzianą na masę, a nie na grubość. W normach IPC, w tabelach możliwości wytwórców i w ustawieniach narzędzi projektowych obowiązuje do dziś, także u wytwórców azjatyckich i europejskich, którzy poza tym posługują się układem metrycznym.

Dwie uwagi praktyczne. Po pierwsze, podana grubość dotyczy folii bazowej; miedź w otworach i na warstwach zewnętrznych po galwanizacji jest grubsza, a jej rzeczywista grubość zależy od procesu. Po drugie, przejście z 1 oz na 2 oz zmienia nie tylko przekrój ścieżki, lecz także minimalną szerokość i odstęp, jakie wytwórca jest w stanie wykonać – grubszej warstwy nie da się wytrawić równie precyzyjnie.

- ☐ Archiwum wydania zapisane z datą i sumą kontrolną.

Czego i gdzie nauczyć się dalej

Dokumenty, do których warto sięgnąć, w kolejności przydatności dla konstruktora zaczynającego pracę seryjną:

- IPC-2221C, Generic Standard on Printed Board Design – podstawa całej rodziny IPC-2220, odstępy izolacyjne, wymagania ogólne.
- IPC-2152, Standard for Determining Current Carrying Capacity in Printed Board Design

– obciążalność prądowa oparta na współczesnych badaniach.

- IPC-7352, Generic Guideline for Land Pattern Design, wraz z IPC-7351B – pola lutownicze.
- IPC-A-600, Acceptability of Printed Boards – kryteria odbioru płytek gołych, przydatne przy reklamacjach; wrócimy do niego w części trzeciej.
- The Gerber Layer Format Specification, rewizja 2024.05, dostępna bezpłatnie ze strony Ucamco – jedyne wiążące źródło co do formatu danych.

Dokumenty IPC są odpłatne, ale ich spisy treści są dostępne publicznie i pozwalają ocenić, czy dany tytuł odpowiada na konkretną potrzebę, zanim wyda się pieniądze.

Z pozycji książkowych podstawowy zestaw dla konstruktora obejmuje opracowanie zbiorowe Printed Circuits Handbook pod redakcją Coombsa i Holdena jako encyklopedię procesu, Signal and Power Integrity – Simplified Erica Bogatina jako wprowadzenie do zagadnień części drugiej oraz Electromagnetic Compatibility Engineering Henry’ego Otta jako klasyczne ujęcie mas i ekranowania.

W części drugiej zajmiemy się tym, co w tym odcinku sygnalizowaliśmy kilkakrotnie i konsekwentnie odkładaliśmy: integralnością sygnału i kompatybilnością elektromagnetyczną. Kiedy „szybko” zaczyna znaczyć „linia długa”, jak naprawdę płynie prąd powrotny, co kontroluje wytwórca przy zamówieniu impedancji kontrolowanej, jak zbudować sieć dystrybucji zasilania i ile kosztują badania EMC w polskim laboratorium.

JT



Wydanie EP05/2026 dostępne na
www.UlubionyKiosk.pl

Złe nawyki w projektowaniu PCB – jak je zwalczać

Eric Bogatin – dziekan Teledyne LeCroy Signal Integrity Academy i profesor Uniwersytetu Kolorado w Boulder – przez trzy dekady szkolił inżynierów na całym świecie w zakresie integralności sygnałów i zasilania. Wykład wygłoszony na konferencji AltiumLive łączy jego doświadczenia z pracy z praktykującymi inżynierami z aktualną pracą dydaktyczną i badawczą. Centralnym pytaniem wykładu jest: skąd biorą się złe nawyki projektowe i dlaczego tak trudno je wykorzenić?



Dwanaście rozdziałów – od oscyloskopu 12-bit po własny park pomiarowy

Między konstruktorem a rzeczywistością zawsze stoi przyrząd. Dobry pomiar prowadzi wprost do przyczyny problemu; zły – kłamie i wysyła w ślepą uliczkę. Ten kurs jest właśnie o tym: jak dobrać przyrząd do zadania, ustawić go poprawnie i czytać jego wskazania krytycznie, żeby mówić prawdę, a nie to, co chcielibyśmy zobaczyć.

Oś kursu jest instrumentalna – rozdział po rozdziale poznajemy kolejne klasy przyrządów i to, co każdy z nich robi naprawdę. Zaczynamy od oscyloskopu (rozdzielczość i ENOB, wyzwalanie, FFT), przez oscyloskop mieszany i dedykowane analizatory – stanów logicznych oraz widma – po tanie przystawki USB i, osobno, po sondy, które równie łatwo ujawniają problem, co go tworzą. Ostatnią część kursu tworzą trzy case'y integracyjne, w których jeden układ badamy wieloma przyrządami naraz (zasilacz impulsowy, precyzyjny tor analogowy, szybka magistrala cyfrowa), oraz praktyczny przewodnik, jak zbudować własny park pomiarowy.

Każdy rozdział pokazuje jakiś mit albo pułapkę: że 12 bitów znaczy 12 efektywnych bitów, że „ładny przebieg” to brak problemu, że sonda jest przezroczysta dla pomiaru. Opieramy się wyłącznie na źródłach klasy inżynierskiej – notach aplikacyjnych i klasycznych primerach producentów (Keysight, Tektronix, Rohde & Schwarz i innych), z udokumentowanymi procedurami i wynikami. Pokazujemy wielu dostawców, w tym również niedrogie konstrukcje azjatyckie.

Kurs składa się z dwunastu rozdziałów i ukazuje się w sześciu kolejnych wydaniach EP, po dwa rozdziały w każdym. W ramce pełna mapa kursu.

Mapa kursu

Wydanie I

1. Oscyloskop 12-bit: rozdzielczość, ENOB, High-Res
2. Wyzwalanie i pomiary automatyczne

Wydanie II

3. Oscyloskop jako analizator widma (FFT)
4. Oscyloskop mieszany (MSO): analog i cyfra w jednym oknie

Wydanie III

5. Analizator stanów logicznych: I²C, SPI, CAN
6. Analizator widma na biurku

Wydanie IV

7. Przystawki USB – mobilne laboratorium
8. Sondy i akcesoria: pomiar mówi prawdę albo kłamie

Wydanie V

9. Case: zasilacz SMPS pod ostrzałem
10. Case: precyzyjny tor analogowy

Wydanie VI

11. Case: szybka magistrala – eye diagram i jitter
12. Twój lab 2027: park pomiarowy i wzorcowanie

Po dwunastu rozdziałach będziesz wiedział, po jaki przyrząd sięgnąć dla rozwiązania danego problemu, jak go ustawić, jak nie dać się zwieść jego wskazaniom – i jak złożyć z tego wszystkiego własne, sensownie dobrane laboratorium.

WM

Oscyloskop 12-bit w praktyce konstruktora: ile bitów naprawdę trafia na ekran

Rozdział 1. 12 bitów to nie zawsze 12 bitów – co rozdzielczość i tryb High-Res dają konstruktorowi, a czego nie załatwią

Oscyloskopy o rozdzielczości 12 bitów przestały być sprzętem z górnej półki laboratoriów akredytowanych i zeszły cenowo na biurko przeciętnego konstruktora. W materiałach producentów pojawia się przy nich obietnica „szesnastokrotnie większej rozdzielczości” niż w klasycznych przyrządach 8-bitowych. Brzmi to jak czysty zysk – więcej bitów, więcej szczegółów, lepszy pomiar. Rzecz w tym, że między liczbą bitów przetwornika a liczbą bitów, które faktycznie trafiają na ekran jako użyteczna informacja, potrafi być spora różnica. Ten rozdział pokazuje, skąd ta różnica się bierze, kiedy dwanaście bitów naprawdę widać w codziennej pracy, a kiedy jest to tylko liczba w karcie katalogowej.

1. Rozdzielczość pionowa: od poziomów kwantyzacji do miliwoltów

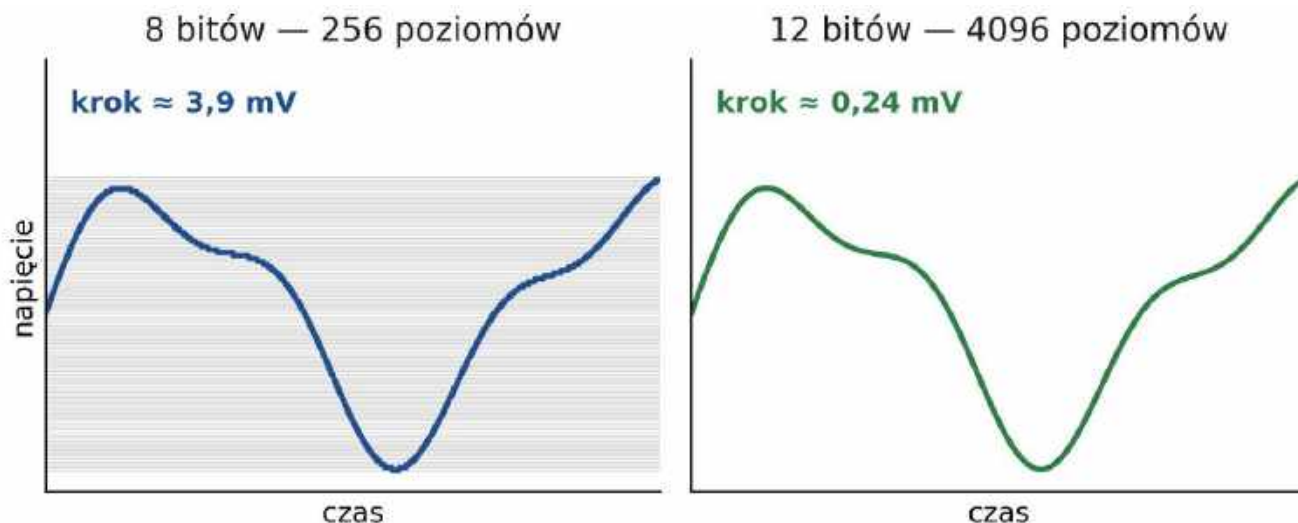
Każdy oscyloskop cyfrowy zamienia napięcie na liczbę w przetworniku analogowo-cyfrowym (ADC). Rozdzielczość pionowa mówi, jak drobne są kwanty, czyli na ile poziomów kwantyzacji przetwornik dzieli zakres wejściowy. Zależność jest wykładnicza: liczba poziomów to dwa do potęgi równej liczbie bitów. Przetwornik 8-bitowy dzieli zakres na 2 do potęgi 8, czyli na 256 poziomów, a 12-bitowy – na 2 do potęgi 12, czyli 4096. Stąd owo „szesnaście razy”: 4096 to dokładnie szesnastokrotność 256.

Sama liczba poziomów niewiele mówi, dopóki nie przełożymy jej na napięcie. Weźmy zakres pełnej skali 1 V. Oscyloskop 8-bitowy dzieli go na 256 kwantów, więc najmniejsza rozróżnialna zmiana to około 3,9 mV. Ten sam zakres w przyrządzie 12-bitowym dzieli się na 4096 kwantów – najmniejsza rozróżnialna zmiana maleje do około 0,24 mV. Różnica robi się namacalna, gdy na szynie zasilania trzeba dojrzeć małe tętnienie. Jeśli na szynie 5 V szukamy tętnienia rzędu kilkunastu miliwoltów, to na przyrządzie 8-bitowym, ustawionym tak, aby zmieścić całe 5 V, tętnienie ma szerokość zaledwie kilku kroków kwantyzacji – ginie w schodkach

albo zamienia się w grubą, poszarpaną kreskę. Przyrząd 12-bitowy widzi w tym samym miejscu gładki przebieg, na którym da się zmierzyć wartość międzyszczytową i policzyć widmo.

Z tym wiąże się pojęcie zakresu dynamicznego – stosunku największego sygnału, który przyrząd obejmuje, do najmniejszego, który jeszcze rozróżnia. Właśnie sytuacja, w której na jednym ekranie musi się zmieścić duży sygnał i jednocześnie mały detal na jego tle, jest naturalnym terenem dla wysokiej rozdzielczości. Im bardziej zgrubny zakres pionowy ustawimy, tym większy krok kwantyzacji w miliwoltach – dlatego problem nie jest abstrakcyjny, lecz pojawia się dokładnie wtedy, gdy nie chcemy albo nie możemy zejść na czulszą skalę.

Formalnie najmniejszy krok nazywa się **LSB** (*least significant bit*), a błąd kwantyzacji pojedynczej próbki mieści się w granicach połowy tego kroku. Błąd rośnie proporcjonalnie do ustawionego zakresu. Przy pełnej skali 1 V krok 8-bitowy to wspomniane 3,9 mV; gdy jednak ustawimy zakres kilku woltów, krok urośnie do kilkunastu czy nawet kilkudziesięciu miliwoltów – i wtedy nawet spore tętnienie potrafi zniknąć między dwoma sąsiednimi poziomami. To dlatego rozdzielczość ujawnia swój brak najostrzej tam, gdzie mały sygnał trzeba



Rysunek 1. Ten sam przebieg na zakresie 1 V: 8 bitów daje krok $\approx 3,9$ mV i widoczne schodki, 12 bitów – krok $\approx 0,24$ mV i przebieg praktycznie gładki

oglądać na tle dużego, a nie tam, gdzie sygnał wypełnia cały ekran.

Warto od razu wprowadzić rozróżnienie, do którego wrócimy w dalszej części: **rozdzielczość to nie to samo co dokładność**. Rozdzielczość mówi, jak drobny jest krok, a nie jak blisko prawdy jest odczyt. Drobny krok jest warunkiem koniecznym, ale nie wystarczającym dokładnego pomiaru małych sygnałów.

2. ENOB – dlaczego 12 bitów to nie 12 bitów

Gdyby przetwornik był idealny, dwanaście bitów oznaczałoby dwanaście bitów czystej informacji. W realnym oscyloskopie tak nie jest, bo na tor pomiarowy nakłada się szereg zakłóceń. Miara, która to opisuje,

jest efektywna liczba bitów – **ENOB** (*effective number of bits*). ENOB mówi, ile bitów rzeczywiście niesie użyteczną informację po odjęciu wszystkiego, co ją psuje: szumu kwantyzacji samego przetwornika, jego nieliniowości różniczkowej (DNL) i całkowitej (INL), szumu termicznego i śrutowego oraz zniekształceń wnoszonych przez wzmacniacz wejściowy.

Warto na chwilę rozłożyć te źródła na czynniki. Szum kwantyzacji to nieusuwalny błąd zaokrąglenia próbki do najbliższego poziomu – jedyny, który wynika wprost z liczby bitów. Nieliniowość różniczkowa (DNL) opisuje, o ile rzeczywiste odstępów kolejnych poziomów odbiegają od ideału, a całkowita (INL) – o ile cała charakterystyka przetwornika odchyła się od linii prostej. Szum termiczny rodzi się w rezystancjach toru wejściowego, śrutowy – w prądach



$$\text{SNR} \approx 6,02 \cdot N + 1,76 \text{ dB}$$

$$12 \text{ bitów} \approx 74 \text{ dB} \quad \cdot \quad 10 \text{ bitów} \approx 62 \text{ dB}$$

SYMULACJA — model poglądowy

Rysunek 2. Z dwunastu bitów nominalnych realnie użyteczne pozostaje około dziesięć; dolne dwa pochłania szum wejścia. Skala SNR wg zależności dla idealnego przetwornika

półprzewodników, a zniekształcenia dokłada wzmacniacz wejściowy, zwłaszcza przy większych amplitudach. Suma tych składników sprawia, że rozdzielczość deklarowana i rozdzielczość użyteczna to dwie różne liczby.

Skalę zjawiska dobrze pokazują pomiary porównawcze publikowane przez producentów i prasę branżową. Dla oscyloskopów z 12-bitowym przetwornikiem realnie wykorzystywane bywa około dziesięciu bitów – dolne dwa bity pochłania szum toru wejściowego, i to w całym zakresie nastaw pionowych. Innymi słowy, przyrząd „dwunastobitowy” pracuje jak dziesięciobitowy, a bywa, że jak dziewięciobitowy: producenci sprzętu z górnej półki podają dla takich konstrukcji stosunek sygnału do szumu rzędu 55 dB.

Skąd te 55 dB i jak je czytać? Dla idealnego przetwornika obowiązuje przybliżenie $SNR \approx 6,02 \times N + 1,76$ dB, gdzie N to liczba bitów. Idealny przetwornik 12-bitowy dałby więc około 74 dB, a dziesięciobitowy – około 62 dB. Deklarowane w praktyce 55 dB odpowiada mniej więcej dziewięciu bitom efektywnym. Nie jest to wada konkretnego przyrządu, lecz naturalna konsekwencja fizyki toru wejściowego – i dlatego liczbę bitów z nagłówka karty katalogowej trzeba czytać z poprawką na szum.

Praktyczny wniosek dla konstruktora: jeśli producent podaje ENOB albo SNR, jest to informacja bardziej rzetelna niż sama liczba bitów przetwornika. A jeśli podaje wyłącznie liczbę bitów – bezpiecznie założyć, że użytecznych będzie o dwa mniej.

3. High-Res: jak oscyloskop 8-bit „dorabia” bity

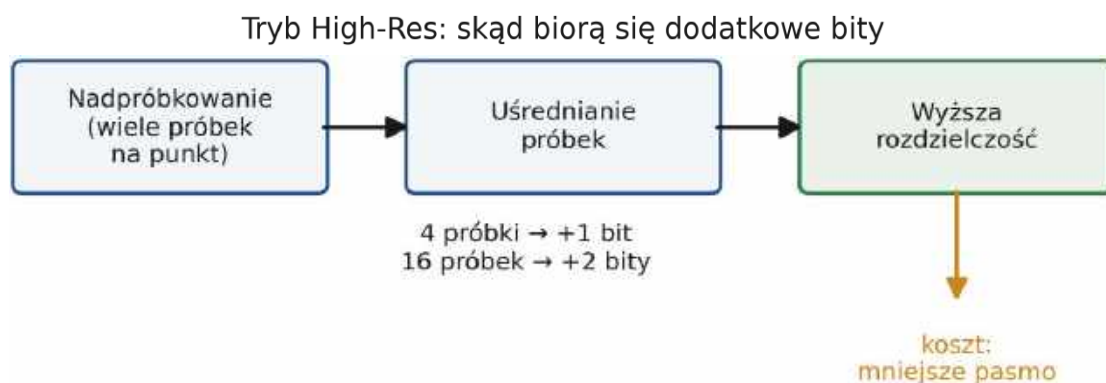
Nie każdy przyrząd, który pokazuje więcej niż osiem bitów, ma dwunastobitowy przetwornik. Wiele oscyloskopów oferuje tryb wysokiej rozdzielczości (**High-Res**), w którym rozdzielczość powstaje programowo, przez **nadpróbkowanie** i **uśrednianie**. Mechanizm jest prosty: jeśli przetwornik próbkuje znacznie szybciej, niż wymaga tego oglądany sygnał, to na jeden punkt wyświetlanego przebiegu przypada wiele próbek. Uśrednienie tych próbek redukuje szum i podnosi efektywną rozdzielczość.

Zysk jest przewidywalny. Uśrednienie czterech kolejnych próbek dokłada mniej więcej jeden bit – z ośmiu robi się dziewięć. Uśrednienie szesnastu próbek dokłada dwa bity – efektywnie dziesięć. Cena też jest przewidywalna i trzeba ją znać: uśrednianie działa jak filtr dolno-przepustowy, więc za każdy dorobiony bit płaci się pasmem. Tryb High-Res świetnie nadaje się do wolnozmiennych sygnałów analogowych i pomiarów małych napięć, ale nie da się nim jednocześnie oglądać szybkich zboczy.

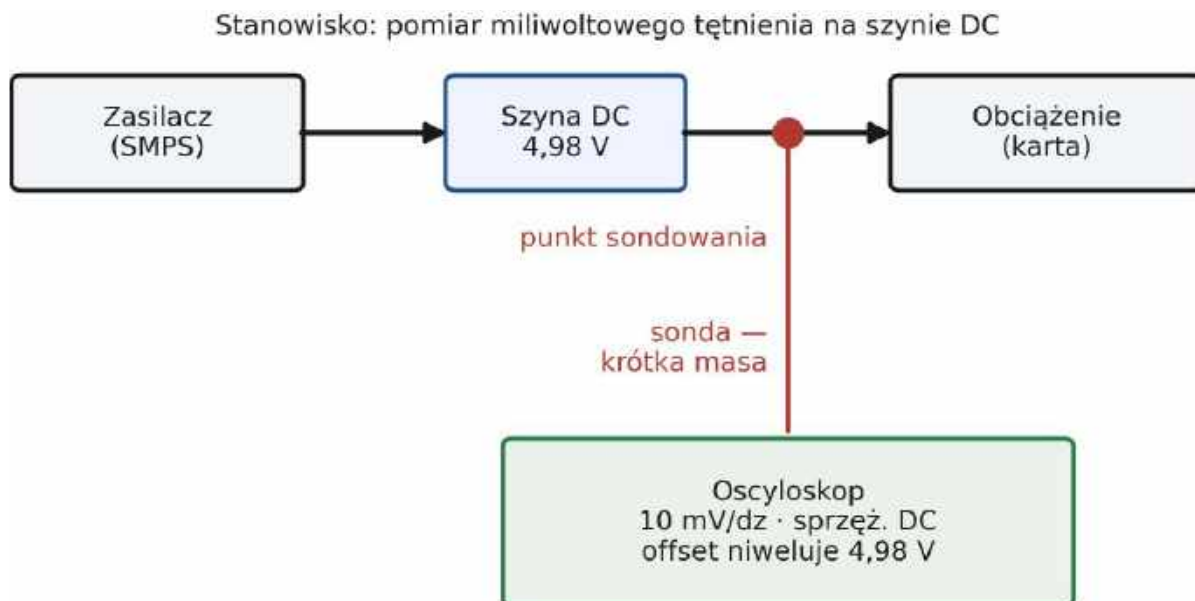
Warto odróżnić High-Res od dwóch pokrewnych trybów akwizycji. Uśrednianie (Average) też redukuje szum, ale wymaga sygnału powtarzalnego i uśrednia po kolejnych wyzwoleniach; High-Res tworzy punkt o wyższej rozdzielczości z jednego przebiegu. Detekcja szczytowa (Peak Detect) robi rzecz przeciwną – zamiast wygładzać, wyłapuje najkrótsze impulsy. Dobór trybu to zawsze świadoma decyzja, co poświęcamy: szum, pasmo czy zdolność łapania wąskich zdarzeń.

To prowadzi do rozróżnienia, które przy zakupie bywa mylące. Sprzętowy przetwornik 12-bitowy daje dwaście bitów „cały czas”, niezależnie od podstawy czasu. Rozwiązania oparte na trybie wysokiej rozdzielczości dorabiają bity programowo – bardzo skutecznie tam, gdzie sygnał jest wolny, ale kosztem pasma tam, gdzie jest szybki. Oba podejścia są w porządku, o ile wiadomo, które się kupuje i do czego. Producenci przyrządów z natywnym 12-bitowym przetwornikiem podkreślają właśnie ten aspekt: pełną rozdzielczość bez obniżania pasma ani częstotliwości próbkowania.

Na rynku spotyka się deklaracje idące dalej – nie tylko dziesięciu czy dziesięciu, ale nawet szesnastu bitów uzyskiwanych programowo. Nie ma w tym oszustwa, dopóki pamiętamy o dwóch rzeczach. Po pierwsze, dorobione bity są warte tyle, ile pozwala na to szum toru wejściowego: jeśli tor szumi na poziomie ograniczającym do dziesięciu bitów, żadne uśrednianie nie wyczaruje szesnastu bitów prawdziwej informacji. Po drugie, każdy dorobiony bit kosztuje pasmo. Dlatego przy takich liczbach pierwsze pytanie brzmi: przy jakim paśmie i przy jakim poziomie szumu wejścia ta rozdzielczość obowiązuje.



Rysunek 3 Tryb High-Res: nadpróbkowanie i uśrednianie zamieniają nadmiar próbek na dodatkowe bity rozdzielczości – kosztem pasma



Rysunek 4 Pomiar miliwoltowego tętnienia na szynie DC: punkt sondowania blisko obciążenia, offset niwelujący napięcie szyny, czuła skala, krótka masa sondy. Oscylogram wyniku – źródło: DigiKey (Art Pini), Fig. 3: www.digikey.com/en/articles/fundamentals-of-8-bit-versus-12-bit-oscilloscopes

4. Kiedy to naprawdę widać – przykład

Najlepiej pokazać różnicę na konkretnym przykładzie.

Przykład dotyczy integralności zasilania i pochodzi z materiału opisującego pomiar na oscyloskopie 12-bitowym klasy WaveSurfer. Na szynie o napięciu 4,98 V zmierzono tętnienie o wartości międzyszczytowej 45 mV, a dolny przebieg stanowiła transformata FFT tego tętnienia, ujawniająca bogate w harmoniczne widmo z podstawową składową około 982 Hz. Kluczowy szczegół procedury: pomiaru dokonano na czułej skali 10 mV/dz, co było możliwe dzięki dużemu zakresowi offsetu – na tej skali przyrząd dysponował offsetem rzędu ± 8 V, więc dało się zniwelować całe 4,98 V szyny i „podnieść” miliwoltowe tętnienie na czułą skalę.

Dlaczego to samo zadanie jest trudne dla przyrządu 8-bitowego? Wróćmy do arytmetyki z pierwszego rozdziału. Na zakresie rzędu 1 V dynamika 256:1 daje najmniejszy krok około 3,9 mV. Tętnienie 45 mV międzyszczytowo to raptem kilkanaście kroków kwantyzacji – a jego drobniejsza struktura, ta, którą ujawnia FFT, tonie w schodkach. Dwanaście bitów, z krokiem 0,24 mV, rozdziela to samo tętnienie na setki poziomów, dając przebieg gotowy do pomiaru wartości międzyszczytowej i analizy widmowej.

Ten przypadek pokazuje też, że rozdzielczość nie działa w próżni. Gdyby przyrząd 12-bitowy nie miał odpowiednio dużego zakresu offsetu, nie dałoby się zejść na skalę 10 mV/dz z całą szyną na wejściu – i cała przewaga rozdzielczości poszłaby na marne. Dlatego w karcie katalogowej obok liczby bitów trzeba czytać także minimalną **czułość** i zakres offsetu.

Jest w tym przykładzie jeszcze jedna lekcja. To, że z jednego zapisu dało się policzyć czyste FFT tętnienia, nie jest przypadkiem: im mniej szumu kwantyzacji w przebiegu

czasowym, tym niższy poziom artefaktów w widmie i tym łatwiej odróżnić prawdziwe harmoniczne od śmieci. Wysoka rozdzielczość pracuje więc podwójnie – raz w dziedzinie czasu, dając gładki przebieg, i drugi raz w dziedzinie częstotliwości, dając czystsze widmo. To pomost do rozdziału trzeciego, w którym oscyloskop posłuży jako „analyzer widma light”, a FFT tętnienia i zakłóceń rozbierzemy na czynniki.

5. Czego 12 bitów nie załatwia

Warto określić granice. **Wyższa rozdzielczość bez odpowiednio niskiego poziomu szumów własnych wejścia jest bezużyteczna** – jeśli dolne bity i tak giną w szumie, ich formalna obecność niczego nie zmienia. Właśnie dlatego producenci przyrządów 12-bitowych tyle uwagi poświęcają konstrukcji toru wejściowego: niskoszumnemu wzmacniaczowi i liniowości stopnia wejściowego. Rozdzielczość jest tylko jednym z trzech filarów dobrego pomiaru małych sygnałów; pozostałe dwa to szum toru i – o czym będzie osobny rozdział – sonda.

Druga granica dotyczy pasma. Dwanaście bitów nie zastąpi pasma tam, gdzie liczą się szybkie zbocza. Co więcej, w konstrukcjach z natywnym 12-bitowym przetwornikiem duża rozdzielczość i szerokie pasmo bywają wzajemnie ograniczane przez architekturę.

Przy porównywaniu przyrządów wysokiej rozdzielczości łatwo wpaść w pułapkę patrzenia na maksymalne wartości każdej specyfikacji z osobna. Tymczasem liczy się to, co przyrząd oferuje jednocześnie: bywa, że deklarowaną najwyższą częstotliwość próbkowania dostaje się tylko na dwóch kanałach, a przy czterech spada ona o połowę. Pomocna jest tu metoda wykresu radarowego, w której na czterech osiach – pasmo, częstotliwość próbkowania,



Rysunek 5 To samo miliwoltowe tętnienie: przez 12 bitów czytelne i gotowe do pomiaru, przez 8 bitów (na zgrubnym zakresie) tonące w schodkach kwantyzacji

rozdzielczość i liczba kanałów – nanosi się parametry dla konkretnej konfiguracji; im większe pole figury, tym lepszy ogólny kompromis. Ta sama metoda pozwala zestawić przyrząd 12-bitowy o mniejszym paśmie z 8-bitowym o większym i wybrać ten, którego pole lepiej pokrywa realne potrzeby. Do tej metody i do doboru całego parku pomiarowego wrócimy w rozdziale dwunastym tego kursu.

Jest też cichy bohater, o którym łatwo zapomnieć: czułość wejścia. To ona decyduje, jak małe napięcie w ogóle da się rozciągnąć na ekranie. Tu przewaga dobrze zaprojektowanych przyrządów wysokiej rozdzielczości jest namacalna – najczulsze zakresy schodzą do 1 mV/dz, a w konstrukcjach z wyjątkowo niskoszumnym wejściem nawet do 500 μ V/dz. Bez takich zakresów sama liczba bitów niewiele daje, bo mały sygnał pozostaje małym sygnałem na zbyt zgrubnej skali.

6. Checklista konstruktora

Zbierzmy to, co przy oscyloskopie wysokiej rozdzielczości warto sprawdzić, zanim uwierzy się nagłówkowej liczbie bitów:

Klasa przyrządu	Najczulszy zakres
Typowy oscyloskop 8-bitowy	zwykle 1...2 mV/dz
Oscyloskop 12-bit (np. Teledyne LeCroy, Agilent)	1 mV/dz
Konstrukcja o wyjątkowo niskim szumie wejścia (np. Rigol)	500 μ V/dz
Przykładowe najczulsze zakresy pionowe wg klasy przyrządu	

Trzy liczby warte zapamiętania

256 kontra 4096 poziomów – 12 bitów to szesnastokrotnie drobniejszy krok niż 8 bitów.

Około 10 – tyle bitów bywa realnie użytecznych z 12 nominalnych; resztę pochłania szum wejścia.

6,02 dB na bit – tyle SNR dokłada każdy bit rozdzielczości idealnego przetwornika.

Kiedy więc sięgać po dwanaście bitów, a kiedy wystarczy osiem? Reguła kciuka jest prosta: im więcej pracy z sygnałami analogowymi, integralnością zasilania, czujnikami i precyzją, tym większy realny zysk z rozdzielczości. Jeśli natomiast oscyloskop służy głównie do oglądania sygnałów cyfrowych i szybkich magistral, priorytetem stają się pasmo i częstotliwość próbkowania, a osiem bitów zwykle wystarcza. Dwanaście bitów to nie magia – to konkretne, mierzalne dziesiąte części miliwolta rozdzielczości w miejscu, w którym przyrząd 8-bitowy pokazuje schodki. Trzeba tylko wiedzieć, że tych bitów jest w praktyce około dziesięciu i że pracują dobrze wyłącznie w towarzystwie niskiego szumu, czułych zakresów i porządnego offsetu.

W następnym rozdziale: wyzwalamie i pomiary automatyczne – ukryta inteligencja oscyloskopu, czyli jak zamienić przyrząd w asystenta, który sam łapie rzadkie zdarzenia.

WM

Źródła:

- Keysight, Evaluating Oscilloscopes oraz glosariusz rozdzielczości pionowej i trybu High-Res.
- Electronic Design/Agilent – materiały o ENOB, SNR i realnie wykorzystywanej liczbie bitów w oscyloskopach 12-bitowych.
- Teledyne LeCroy – materiały o technologii HD4096 oraz metoda porównywania oscyloskopów wysokiej rozdzielczości (radar chart).
- DigiKey/Teledyne LeCroy, Fundamentals of 8-Bit Versus 12-Bit Oscilloscopes – udokumentowany pomiar tętnienia 45 mV międzyszczytowo i FFT ze składową 982 Hz na WaveSurfer 4104HD.

Rozdział 2. Wyzwalanie i pomiary automatyczne: ukryta inteligencja oscyloskopu

Jak zamienić oscyloskop w asystenta, który sam łapie rzadkie zdarzenia – i jak ufać temu, co pokazują pomiary automatyczne.

Pasmo i częstotliwość próbkowania decydują o tym, czy zdarzenie da się w ogóle zobaczyć. Wyzwalanie decyduje o czymś równie ważnym: czy uda się je złapać. W typowej pracy konstruktora układ działa poprawnie przez 99,999% czasu, a problem kryje się w tym ułamku procenta – w rzadkim **glitchu**, w jednej ramce na magistrali, w pojedynczym naruszeniu czasowym. Najlepszy oscyloskop świata, ustawiony na zwykłe **wyzwalanie zboczem**, może takie zdarzenie przeoczyć albo utopić w gąszczu nieistotnych przebiegów. Ten rozdział jest o tym, jak zamienić oscyloskop z biernego ekranu w asystenta, który sam wyławia interesujące zdarzenia – oraz o drugiej stronie tej automatyzacji: pomiarach automatycznych i o tym, kiedy można im ufać.

1. Po co wyzwalanie

Wyzwalanie (trigger) synchronizuje poziomą podstawę czasu z sygnałem – mówi oscyloskopowi, w którym momencie zacząć rysować przebieg. Bez niego kolejne zapisy zaczynałyby się w przypadkowych miejscach sygnału i nakładały na siebie w nieczytelną plątaninę. Z wyzwalaniem powtarzalne przebiegi „zastygają” w jednym miejscu ekranu, a pojedyncze zdarzenia dają się uchwycić i zatrzymać. To najstarsza i najważniejsza funkcja oscyloskopu, a mimo to wciąż najczęściej używana w najprostszej postaci.

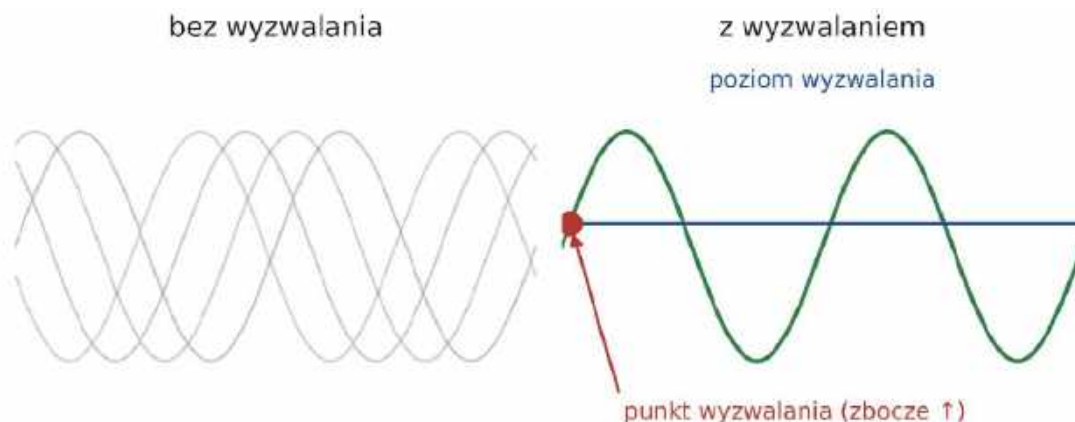
Podstawą jest wyzwalanie zboczem (edge). Definiuje się dwie rzeczy: poziom napięcia, przez który sygnał ma przejść, oraz kierunek – zbocze narastające albo opadające.

Punkt, w którym sygnał przecina ten poziom w zadanym kierunku, staje się punktem odniesienia dla całego zapisu. Jeśli poziom ustawimy powyżej dodatniego szczytu sygnału lub poniżej ujemnego, wyzwalanie zniknie, a obraz znów się rozjedzie.

Samo wyzwalanie zboczem zwykle wystarcza, by zobaczyć podstawową amplitudę i czas przebiegu, dlatego pozostaje domyślnym wyborem. Gdy jednak sygnał jest zaszumiony, pomagają dodatkowe ustawienia toru wyzwalania: sprzężenie (AC albo DC) oraz filtry tłumienia szumu i wysokich lub niskich częstotliwości, które zapobiegają fałszywym wyzwoleniom, nie zmieniając samego sygnału na ekranie.

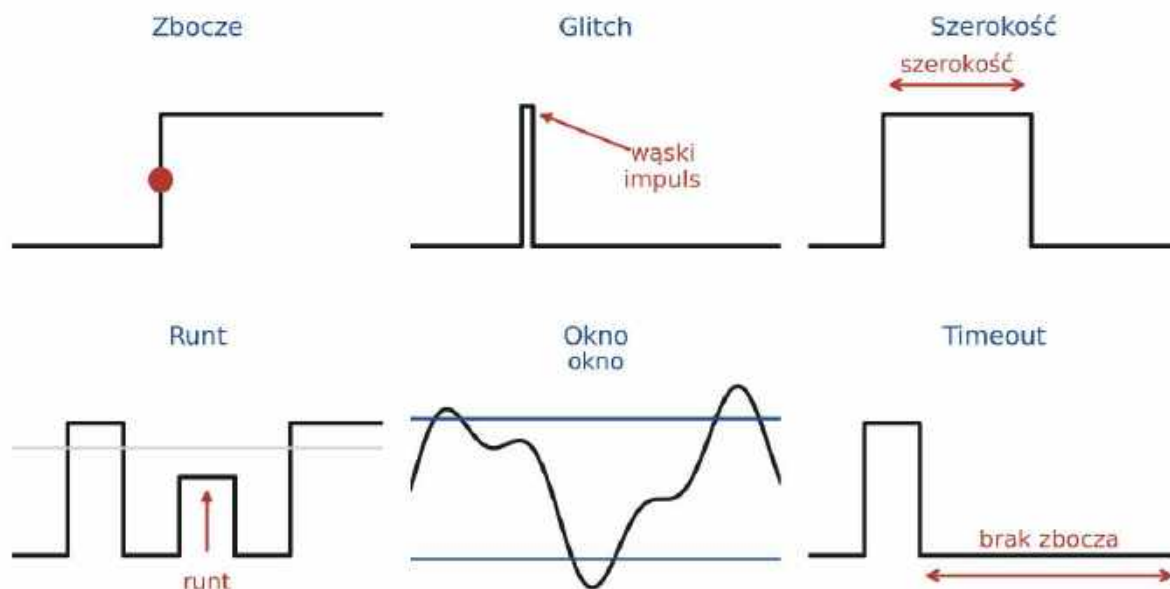
Do tego dochodzą tryby i kwalifikatory, które warto rozumieć. Tryb Normal rysuje przebieg tylko wtedy, gdy pojawia się prawdziwe zdarzenie wyzwalające – idealny do polowania na rzadkie zdarzenia, bo ekran nie „ożywa” bez powodu. Tryb Auto dorysowuje przebieg także wtedy, gdy zdarzenia nie ma, co bywa wygodne przy wstępnym ustawianiu, ale potrafi ukryć fakt, że właściwy trigger nie zachodzi. Osobnym, często niedocenianym narzędziem jest **holdoff** – czas, przez który po wyzwoleniu oscyloskop celowo ignoruje kolejne zdarzenia, aż upłynie zadany odstęp i pojawi się nowe zbocze. Holdoff porządkuje obraz sygnałów o złożonej, powtarzalnej strukturze, na przykład paczek impulsów, gdzie zwykle wyzwalanie łapałoby raz jedno, raz inne zbocze paczki.

Warto też pamiętać, że oscyloskop nie musi wyzwalać się na tym sygnale, który ogląda. Źródłem wyzwalania może być dowolny kanał – także niewyświetlany – albo osobne wejście



Rysunek 1 Bez wyzwalania kolejne zapisy nakładają się w nieczytelną plątaninę; z wyzwalaniem przebieg zastyga względem poziomu wyzwalania i punktu przecięcia na zboczu narastającym

Typy wyzwalania zaawansowanego



Rysunek 2. Sześć typowych warunków wyzwalania zaawansowanego. Każdy odpowiada na inne pytanie diagnostyczne – od krótkiego zakłócenia (glitch) po brak spodziewanego zbocza (timeout).

wyzwalania, a większość przyrządów udostępnia również wyjście wyzwalania, którym można zsynchronizować inny przyrząd: generator, licznik czy drugi oscyloskop. Ta swoboda bywa kluczowa, gdy interesujące zdarzenie widać wyraźnie na jednej linii, na przykład na sygnale zezwolenia, a analizować chcemy inną.

Holdoff najlepiej zrozumieć na przykładzie paczki impulsów powtarzanej co pewien czas. Bez holdoffu oscyloskop wyzwoli się na pierwszym napotkanym zboczu w paczce – raz na trzecim, raz na piątym impulsie – i obraz będzie „skakał”. Ustawienie holdoffu nieco dłuższego niż czas trwania paczki sprawia, że przyrząd ignoruje zbocza we wnętrzu paczki i czeka na początek następnej, dając stabilny obraz. To prosty parametr, a rozwiązuje problem, który potrafi latami frustrować przy złożonych przebiegach.

2. Poza zboczem – wyzwalanie zaawansowane

Wyzwalanie zboczem pokazuje amplitudę i czas, ale nie odróżnia sytuacji, które dla konstruktora są kluczowe: krótkiego zakłócenia od poprawnego impulsu, impulsu za wąskiego od za szerokiego, przebiegu, który „nie dojechał” do właściwego poziomu. Dlatego nowoczesne oscyloskopy oferują całą rodzinę wyzwalających zaawansowanych, a każdy z nich odpowiada na inne pytanie diagnostyczne.

Najważniejsze typy to: glitch – wyłapuje albo odrzuca impulsy krótsze od zadanej granicy, co pozwala tropić przypadkowe zakłócenia; szerokość (width) – wyzwala na impulsach węższych lub szerszych niż zadana wartość; runt – łapie impulsy, które przekroczyły jeden próg, ale nie osiągnęły drugiego, co jest typowym objawem

problemów z poziomami logicznymi; okno (window) – reaguje, gdy sygnał opuszcza zadany pas napięć, w górę albo w dół; timeout – wyzwala, gdy przez określony czas nie pojawi się spodziewane zbocze, co świetnie nadaje się do wykrywania „zawieszeń”; setup&hold – łapie naruszenia czasu ustalania i podtrzymania między danymi a zegarem; wreszcie pattern i state – reagują na zadaną kombinację stanów na kilku liniach.

Skalę możliwości dobrze pokazują liczby. Rozbudowany system wyzwalania potrafi wykrywać glitche o szerokości do 150 ps przy minimalnym czasie ponownego uzbrojenia rzędu 300 ps. To nie jest już „rysowanie sygnału” – to precyzyjne sito zdarzeń, które przepuszcza wyłącznie to, czego szukamy.

Do tej rodziny należy jeszcze wyzwalanie czasem przejścia (transition), reagujące na zbyt wolne lub zbyt szybkie zbocza, oraz – na przyrządach o najwyższych osiągnięciach – wyzwalanie wzorcem szeregowym. To ostatnie potrafi wyzwolić na zadanej sekwencji bitów o długości nawet 64 bitów (40 bitów dla danych kodowanych 8b/10b) i przy strumieniach do 6,25 Gb/s, z zegarem osobnym albo wbudowanym; w praktyce pozwala złapać konkretny fragment transmisji na szybkiej magistrali szeregowej, na przykład określony wzorec w strumieniu PCI Express.

Dobór typu do usterki jest w praktyce prosty. Glitch i szerokość służą do polowania na zakłócenia oraz impulsy o nieprawidłowym czasie trwania; runt zdradza problemy z poziomami logicznymi na styku różnych standardów napięciowych; okno wychwytyuje wyjścia sygnału poza dozwolony zakres, na przykład zapady i skoki zasilania; timeout demaskuje zawieszenia i brakujące zbocza zegara;



Rysunek 3. Wyzwalanie sekwencyjne: zdarzenie A uzbraja układ, a dopiero zdarzenie B – po spełnieniu warunku czasu lub stanu – wyzwala zapis

a setup&hold jest wprost stworzony do łapania naruszeń czasowych na szybkich magistralach równoległych i w pamięciach.

3. Wyzwalanie sekwencyjne A → B

Część zdarzeń da się opisać dopiero w kontekście: „wyzwól na impulsie danych, ale tylko wtedy, gdy wcześniej pojawił się sygnał zezwolenia”. Do tego służy **wyzwalanie sekwencyjne**, w którym definiuje się dwa zdarzenia: zdarzenie A uzbraja układ wyzwalania, a dopiero zdarzenie B – po spełnieniu warunku czasu lub stanu – powoduje właściwe wyzwolenie. W rozbudowanych systemach oba zdarzenia mogą korzystać z pełnej palety wyzwań zaawansowanych, a nie tylko ze zbocza; taką architekturę Tektronix nazywa wyzwalaniem dwuzdarzeniowym A i B.

Dzięki temu można na przykład zażądać: „uzbrój się na zboczu sygnału RESET, a wyzwól na pierwszym impulsie zegara, który pojawi się później niż 2 μs po nim”. Tego rodzaju kwalifikacja zamienia oscyloskop w narzędzie do łapania zdarzeń osadzonych w konkretnej sekwencji, a nie wyrwanych z kontekstu – co w praktyce odróżnia debugowanie „na oślep” od celnego trafienia w przyczynę.

Sekwencyjne wyzwalanie ma naturalne przedłużenie w funkcjach wyszukiwania i znakowania (Search & Mark). Gdy zdarzenie już raz złapiemy, oscyloskop potrafi automatycznie przejrzeć cały długi zapis i oznaczyć wszystkie miejsca spełniające ten sam warunek, a przyciskami „poprzedni” i „następny” skaczemy między nimi. To zamienia pojedyncze trafienie w systematyczny przegląd – bezcenne, gdy trzeba policzyć, jak często dana usterka występuje.

W wielu przyrządach zdarzenie B można dodatkowo opóźnić o zadany czas albo o zadaną liczbę zdarzeń – na przykład „wyzwól na dziesiątym impulsie zegara po zboczu A”. To pozwala celnie wejść w środek długiej sekwencji bez ręcznego przewijania zapisu.

4. Visual Trigger – rysujesz warunek na ekranie

Nawet najlepszy sprzętowy układ wyzwalania rozpoznaje ograniczoną liczbę cech sygnału, i to opisanych

liczbami: poziomami, czasami, szerokościami. Tymczasem oscyloskop jest urządzeniem wizualnym, a wiele warunków najłatwiej zdefiniować, rysując je. Na tym polega **Visual Trigger**: użytkownik zaznacza na ekranie obszary lub kształty, a oscyloskop przegląda kolejne zarejestrowane przebiegi i odrzuca te, które do nich nie pasują. W efekcie na ekranie, w pomiarach i w pamięci referencyjnej zostają wyłącznie przebiegi spełniające narysowany warunek. Różni to Visual Trigger od testu maski: tu przebiegi niepasujące są aktywnie odrzucane, a nie tylko sygnalizowane.

Klasyczny przykład dotyczy pamięci DDR3. Aby wyizolować rzadkie zapisy o określonym wzorcu bitów, ustawia się sprzętowe wyzwalanie na paczce zapisu DQ o wartości 11000000, a następnie dodaje graficzny obszar „keep-in” na sygnale DQ, który przepuszcza zdarzenie tylko wtedy, gdy start następuje z wartości innej niż stan wysokiej impedancji. Tak złapane, rzadkie zdarzenie można potem analizować pod kątem zgodności czasowej. Czytelnika, który chce zobaczyć oryginalne zrzuty, odsyłamy do materiału producenta: Tektronix, Visual Triggering.

W praktyce Visual Trigger zawsze zaczyna się od ustawienia sprzętowego wyzwalania w trybie Normal – może to być zwykłe zbocze albo złożony trigger szeregowy – a dopiero na tak zebrane przebiegi nakłada się warunek graficzny. Ten sam mechanizm da się wykorzystać nie tylko do wyzwalania, ale i do wyszukiwania zdarzeń w już zarejestrowanym materiale oraz do porównania z przebiegiem wzorcowym.

5. Wyzwalanie na magistralach szeregowych

Największą wartość wyzwalanie pokazuje na magistralach szeregowych, gdzie interesujące zdarzenie to konkretna treść: adres, dane, typ ramki, błąd. Problem w tym, że klasyczny trigger stanowy tu nie wystarcza – musiałby mieć głębokość dziesiątek, a nawet setek stanów, po jednym na każdy bit, i trzeba by go mozolnie programować bit po bicie. Dlatego oscyloskopy z opcją analizy magistral wyzwalają się bezpośrednio na treści zdekodowanej ramki.

Na magistrali I²C można na przykład zażądać wyzwolenia na konkretnym adresie i danych; w materiałach producenta pokazano trigger na adresie 0x50 i danych 0x00, z automatycznym zdekodowaniem całej ramki (przy czym adres da się interpretować w schemacie 7-bitowym z osobnym bitem R/W albo 8-bitowym). Na magistrali CAN dostępne są warunki dopasowane do struktury protokołu: wyzwalamie na polu początku ramki (Start of Frame) oraz na typie ramki – do wyboru ramka danych, zdalna, błędna i przeciążenie. To bezcenne przy tropieniu sporadycznych błędów: zamiast wyzwalać na czymkolwiek i ręcznie przeglądać ramkę po ramce, każe się oscyloskopowi czekać właśnie na ramkę błędna. Do samego dekodowania magistral wrócimy szerzej w rozdziale 5.

Warto dodać kilka praktycznych szczegółów. CAN jest magistralą różnicową i choć oscyloskop zdekoduje ją także sondą pojedynczą, wierność sygnału i odporność na zakłócenia rosną przy sondowaniu różnicowym. Zdekodowaną transmisję wygodnie ogląda się w formie tabeli wyników ze znacznikami czasu, którą łatwo zestawzić z listingiem oprogramowania; kliknięcie wiersza tabeli przenosi oscyloskop do odpowiadającego mu fragmentu przebiegu. Analogiczne mechanizmy działają dla SPI, UART, LIN czy FlexRay – na LIN można na przykład wyzwolić na identyfikatorze 01h, po którym następuje wartość danych 1Eh. Bez wyzwalamie na treści trzeba by wyzwolić na czymkolwiek, zarejestrować możliwie długie okno i ręcznie dekodować ramkę po ramce, aż trafi się na właściwą – zadanie, które potrafi zająć od kilkunastu minut do godzin.

Podobnie działają pozostałe magistrale. Na SPI można wyzwolić na konkretnym słowie danych zsynchronizowanym z sygnałem wyboru układu, a na UART – na zadanym bajcie albo na błędzie parzystości. Wspólny mianownik jest ten sam: oscyloskop rozumie strukturę protokołu, więc warunek zapisuje się w kategoriach ramki, a nie napięć i czasów.

6. Pomiary automatyczne – i jak im ufać

Drugą twarzą „inteligentnego” oscyloskopu są pomiary automatyczne. Zamiast odczytywać wartości z podziałki,

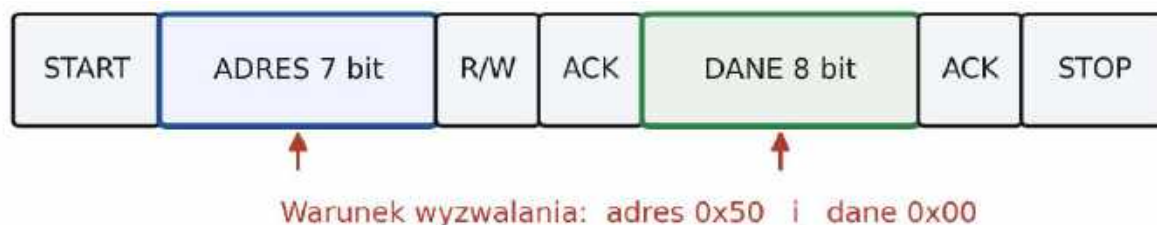
każe się przyrządowi policzyć je samodzielnie: amplitudowe – wartość międzyszczytową (V_{pp}), skuteczną (RMS), średnią; czasowe – okres i częstotliwość, szerokość impulsu, opóźnienie i przesunięcie fazowe. Współczesne oscyloskopy mają też wbudowane przyrządy pomocnicze: cyfrowy woltomierz i częstotłomierz o rozdzielczości sięgającej kilku, a nawet kilkunastu cyfr, którego dokładność zależy wprost od dokładności wzorca podstawy czasu.

Kluczowe jest zrozumienie, jak te pomiary są definiowane, bo od tego zależy ich wiarygodność. Czas narastania mierzy się standardowo między poziomem 10% a 90% amplitudy sygnału, przy czym poziomy 0% (Base) i 100% (Top) wyznacza się z samego przebiegu, a progi bywają konfigurowalne (na przykład 20...80%). Pomiar można też ograniczyć do wybranego fragmentu przebiegu (bramkowanie, measurement region), a jego wielokrotne powtórzenie daje statystykę: minimum, maksimum, średnią i odchylenie – bezcenne przy ocenie stabilności.

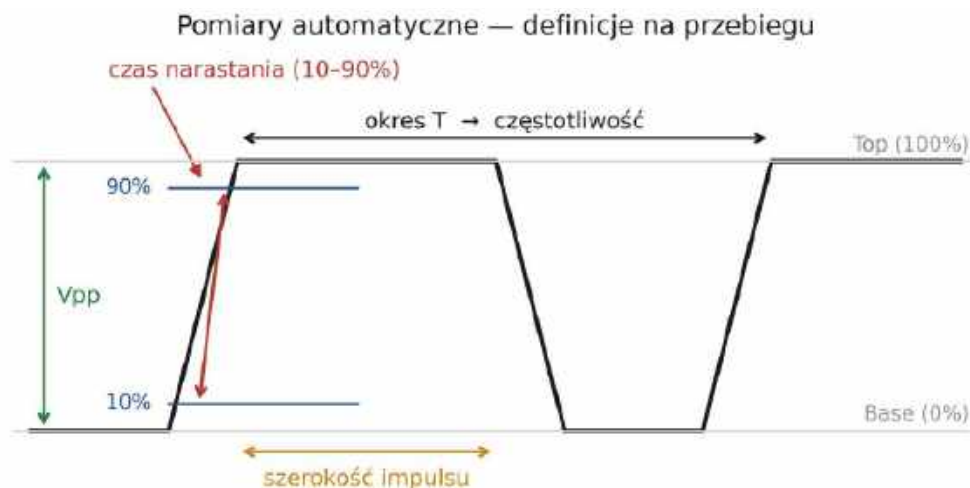
Automatyzacja nie zwalnia jednak z myślenia, bo **pomiary automatyczne** mają swoje pułapki. Najważniejsza dotyczy pasma: jeśli tor pomiarowy jest zbyt wolny, zmierzony **czas narastania** będzie sztucznie wydłużony. Dla toru o charakterystyce gaussowskiej obowiązuje przybliżenie czasu narastania $\approx 0,339/\text{pasma}$, a w praktyce inżynierskiej przyjęło się $0,35/\text{pasma}$; stąd reguła kciuka, by czas narastania toru był około pięć razy krótszy od mierzonego. Druga pułapka to szum, który zawyża wartości międzyszczytowe i rozmywa pomiary czasowe – tu pomaga ograniczenie pasma i uśrednianie. Trzecia to **aliasing** i zbyt mała liczba próbek na okres: częstotliwość policzona z niedoprobkowanego przebiegu potrafi być myląca. Zasada jest prosta: zanim uwierzysz liczbie, spójrz, czy przebieg, z którego powstała, jest poprawnie odwzorowany.

Zakres pomiarów automatycznych jest szeroki. Poza wymienionymi dochodzą czas opadania, współczynnik wypełnienia, przeregulowanie (overshoot) i podszczyt, a przy pomiarach dwukanałowych – opóźnienie i faza między sygnałami. Wbudowany woltomierz oferuje zwykle kilka trybów: międzyszczytowy, skuteczny AC, składową stałą i skuteczny DC, a licznik częstotliwości sięga rozdzielczości od pięciu do nawet dziesięciu cyfr i mierzy do pasma

Wyzwalanie na treści ramki I²C



Rysunek 4. Wyzwalanie na treści ramki I²C: zamiast na surowym zboczu, oscyloskop czeka na zadany adres i dane (tu: adres 0x50 i dane 0x00). Oscylogram dekodowania (zrzut producenta) – źródło: Tektronix, Debugging Serial Buses – oscylogram I²C (rysunek 7) oraz dekodowanie CAN (rysunek 31)



Rysunek 5. Definicje najważniejszych pomiarów automatycznych: wartość międzyszczytowa (V_{pp}), okres i częstotliwość, szerokość impulsu oraz czas narastania między poziomami 10% i 90%

Co to jest aliasing

Mechanizm ilustruje rysunek powyżej: szybki sygnał rzeczywisty (szary), spróbkowany zbyt rzadko (kropki), zostaje odtworzony jako wolny, fałszywy przebieg (niebieski). Twierdzenie o próbkowaniu (Nyquista-Shannona) mówi, że aby wiernie odtworzyć sygnał, trzeba go próbować z częstotliwością co najmniej dwukrotnie wyższą od najwyższej składowej. Gdy ten warunek jest złamany – sygnał ma składowe powyżej połowy częstotliwości próbkowania – oscyloskop „widzi” je jako fałszywe, nieistniejące składowe o niższej częstotliwości. Na ekranie objawia się to nieruchomym, ale zmyślnym przebiegiem albo charakterystycznym „pływaniem” kształtu. Praktyczna reguła: częstotliwość próbkowania powinna z zapasem przewyższać najwyższą częstotliwość w sygnale (w praktyce dąży się do co najmniej kilku próbek na najkrótsze zbocze), a tor o stromej, płaskiej charakterystyce dodatkowo tłumi składowe powyżej częstotliwości Nyquista, zmniejszając ryzyko aliasingu. Jeśli pomiar częstotliwości albo kształt przebiegu wygląda podejrzanie, najpierw zwiększ częstotliwość próbkowania i sprawdź, czy obraz się nie zmienia.

przrządu, korzystając z poziomu wyzwiania jako progu zliczania. Jego dokładność jest tak dobra, jak dokładność wewnętrznego wzorca podstawy czasu – typowo od jednostek do kilkudziesięciu ppm – i można ją poprawić, podając zewnętrzny wzorzec 10 MHz.

Na koniec drobiazg, który łączy pomiary z próbkowaniem: tor o płaskiej charakterystyce mierzy czas narastania dokładniej niż tor gaussowski o tym samym paśmie i skuteczniej tłumi składowe powyżej częstotliwości Nyquista, redukując błędy aliasingu – kosztem



Rysunek 6. Aliasing: gdy na okres sygnału przypadają mniej niż dwie próbki, próbki szybkiego przebiegu (szary) układają się w wolny, nieistniejący przebieg (niebieski) – alias. Dlatego pomiar częstotliwości z niedoprobkowanego sygnału potrafi mocno mylić

Tor gaussowski a płaski

Charakterystyka częstotliwościowa toru wejściowego opisuje, jak oscyloskop przenosi kolejne częstotliwości aż do pasma (punktu -3 dB). Tor gaussowski, typowy dla klasycznych, analogowych konstrukcji z długim łańcuchem wzmacniaczy, opada łagodnie i nie wnosi przeregulowania – jego czas narastania wiąże się z pasmem zależnością ok. $0,339/\text{pasma}$. Tor płaski (brick-wall), spotykany w nowoczesnych oscyloskopach cyfrowych z obróbką DSP, przenosi częstotliwości poniżej pasma równiej, a powyżej tłumi je bardzo stromo. Dzięki temu mierzy czas narastania dokładniej i skuteczniej tłumi aliasing, ale w odpowiedzi na szybki skok potrafi wnieść lekkie przeregulowanie i dzwonienie. W praktyce inżynierskiej dla obu przyjęto się przybliżenie czas narastania $\approx 0,35/\text{pasma}$; przy torze płaskim lepiej polegać na paśmie podanym przez producenta niż na tej regule.

lekkiego przeregulowania w odpowiedzi na skok. To kolejny powód, by liczbie z pomiaru automatycznego zawsze towarzyszyło spojrzenie na sam przebieg.

Dwie rzeczy warto ustawić świadomie. Po pierwsze, poziomy referencyjne: pomiar czasu narastania między 10% a 90% da inną liczbę niż między 20% a 80%, więc porównując wyniki, trzeba trzymać się tej samej definicji. Po drugie, liczba okresów na ekranie: stabilny pomiar okresu i częstotliwości wymaga, by na zapisie mieściło się kilka pełnych cykli i dość próbek na każdy z nich – wartości policzone z jednego, ledwie widocznego okresu potrafią mocno zmylić.

Trzy rzeczy warto zapamiętania

Tryb Normal plus właściwy typ triggera zamieniają oscyloskop w sito zdarzeń.

Na magistrali wyzwalaj na treści, nie na zboczach – trigger stanowy „bit po bicie” to ślepa uliczka.

Pomiar automatyczny jest tak dobry, jak przebieg, z którego powstał: pasmo, szum i próbkowanie decydują o zaufaniu.

7. Checklista konstruktora

W następnym rozdziale: oscyloskop jako analizator widma (FFT) – jak z jednego zapisu wyczytać częstotliwości zakłóceń i tętnień, i gdzie leżą granice tej metody. **WM**

Źródła

- Tektronix, Triggering Fundamentals with Pinpoint Triggering and Event Search & Mark – typy wyzwalania, glitch do 150 ps, ponowne uzbrojenie 300 ps, wyzwalanie dwuzdarzeniowe A/B, wyzwalanie wzorcem szeregowym.
- Tektronix, Visual Triggering – graficzne definiowanie warunku, przykład DDR3, obszar „keep-in”.
- Tektronix, Oscilloscope Systems and Controls – podstawy wyzwalania, holdoff, tryby Auto/Normal, wyzwalanie magistral szeregowych.
- Tektronix, Debugging Serial Buses in Embedded System Designs oraz Debugging CAN, LIN and FlexRay Automotive Buses – I²C 0x50/0x00, typy ramki CAN, dekodowanie.
- Keysight, Understanding Oscilloscope Frequency Response... oraz materiały o pomiarach – czas narastania 10...90%, zależność 0,35/pasmo, reguła 5:1, woltomierz i częstotściomierz w oscyloskopie.

REKLAMA

m.technik
Ciekawi świata są zawsze młodzi

**Ciekawszy,
na czasie,
na topie...
wiecznie młody!**

przejrzyj i kupisz na stronie
www.ulubionykiosk.pl

m.technik
Q-DAY
ZDAŻYĆ PRZED KATASTROFĄ
KIEDY KOMPUTERY KWANTOWE
ZŁAMIA SZYFRY ŚWIATA
science fiction w „Młodym Techniku”
Tomasz Nowak: Entropia

Rozdział 5. MIFARE Ultralight i EV1: pamięć bez Crypto1

MIFARE Classic jest rozbudowany, przy tym skomplikowany – Crypto1, klucze, bity dostępu. MIFARE Ultralight jest jego przeciwieństwem: prosta i tania karta bez żadnej kryptografii, zaprojektowana do zastosowań, w których wystarczy samo przechowywanie małej ilości danych. Ultralight EV1 to nowsza generacja, która dokłada liczniki jednorazowe i opcjonalne hasło, pozostając przy tej samej prostocie interfejsu. Oba typy mają identyczne SAK i ATQA jak NTAG216 – rozróżniamy je dopiero przez komendę GET_VERSION.

Architektura MIFARE Ultralight (MFOICU1)

Ultralight MFOICU1 ma 64 bajty pamięci podzielone na 16 stron po 4 bajty. To podstawowa różnica wobec MIFARE Classic: zamiast 16-bajtowych bloków używamy 4-bajtowych stron numerowanych 0–15. Karta odpowiada na komendę READ, zwracając cztery kolejne strony naraz (16 bajtów łącznie) – w naszym kodzie bierzemy tylko pierwszą czwórkę.

Nie ma żadnej autentykacji kryptograficznej. Każdy czytnik może odczytać i zapisać dowolną stronę danych bez klucza. Jedyną nieodwracalną ochroną

są lock bity w stronach 2 i 3 – po ustawieniu nie można ich cofnąć. W kursie nigdy nie ustawiamy lock bitów.

Strony specjalne

Strona	Zawartość	Bezpieczny zapis?
0	UID[0..2] + BCC0	Nie – fabryczny
1	UID[3..6]	Nie – fabryczny
2	BCC1 + bajt wewn. + Lock[0..1]	Ostrożnie – lock bity!
3	OTP – lock bity jednorazowe	Ostrożnie – nieodwracalne!
4–15	Dane użytkownika (48 B)	Tak – swobodnie

Strona	B0	B1	B2	B3	Zawartość
0	UID0	UID1	UID2	BCC0	Serial number part 1 + checksum
1	UID3	UID4	UID5	UID6	Serial number part 2
2	BCC1	INT	LOCK0	LOCK1	BCC1 + wewn. + lock bity (ostrożnie!)
3	OTP0 – OTP3				One-Time Programmable — nieodwracalne lock bity!
4	bajt[0] – bajt[3]				dane użytkownika
5	bajt[0] – bajt[3]				dane użytkownika
6–14					
15	bajt[0] – bajt[3]				ostatnia strona danych

- Strony 0–1: UID (7 B) i checksuma — fabryczne, tylko odczyt
- Strona 2: lock bity — zapis możliwy, nieodwracalny!
- Strona 3: OTP — nieodwracalne lock bity dla stron danych
- Strony 4–15: 48 B danych użytkownika — swobodny zapis

1. Mapa pamięci MIFARE Ultralight MFOICU1 (16 stron × 4 B = 64 B). Strony 0–1 = UID, strona 2 = lock bity, strona 3 = OTP (nieodwracalne), strony 4–15 = 48 B danych użytkownika. Źródło: opracowanie własne na podstawie: NXP MFOICU1 MIFARE Ultralight Datasheet, Rev. 3.4 – nxp.com

Lock bity są nieodwracalne

Ustawienie bitu lock na 1 blokuje odpowiednią stronę przed zapisem na zawsze. W kursie używamy wyłącznie stron 4–15 – klasa `write_page()` podaje wyjątek przy próbie zapisu do stron 0–3.

MIFARE Ultralight EV1 – co nowego

EV1 dostępny jest w dwóch rozmiarach: MF0UL11 (48 B danych użytkownika, 20 stron łącznie) i MF0UL21 (128 B danych, 41 stron). Oba mają identyczne SAK=0x00 jak klasyczny Ultralight – jedyną pewną metodą odróżnienia jest polecenie `GET_VERSION` (0x60), które klasyczny MF0ICU1 odrzuca jako błąd.

GET_VERSION – jak odróżnić Ultralight od EV1 oraz od NTAG

`GET_VERSION` zwraca 8 bajtów. Kluczowy jest bajt `product_type`: 0x03 oznacza Ultralight EV1, 0x04 oznacza NTAG21x (rozdziały 6...7). Klasyczny MF0ICU1

nie zna tej komendy – zwraca błąd. Nasza klasa `MifareUltralightEV1` korzysta z tej logiki: `main5.py` próbuje `GET_VERSION` w bloku `try/except` i na tej podstawie wybiera właściwą klasę.

Liczniki 24-bitowe

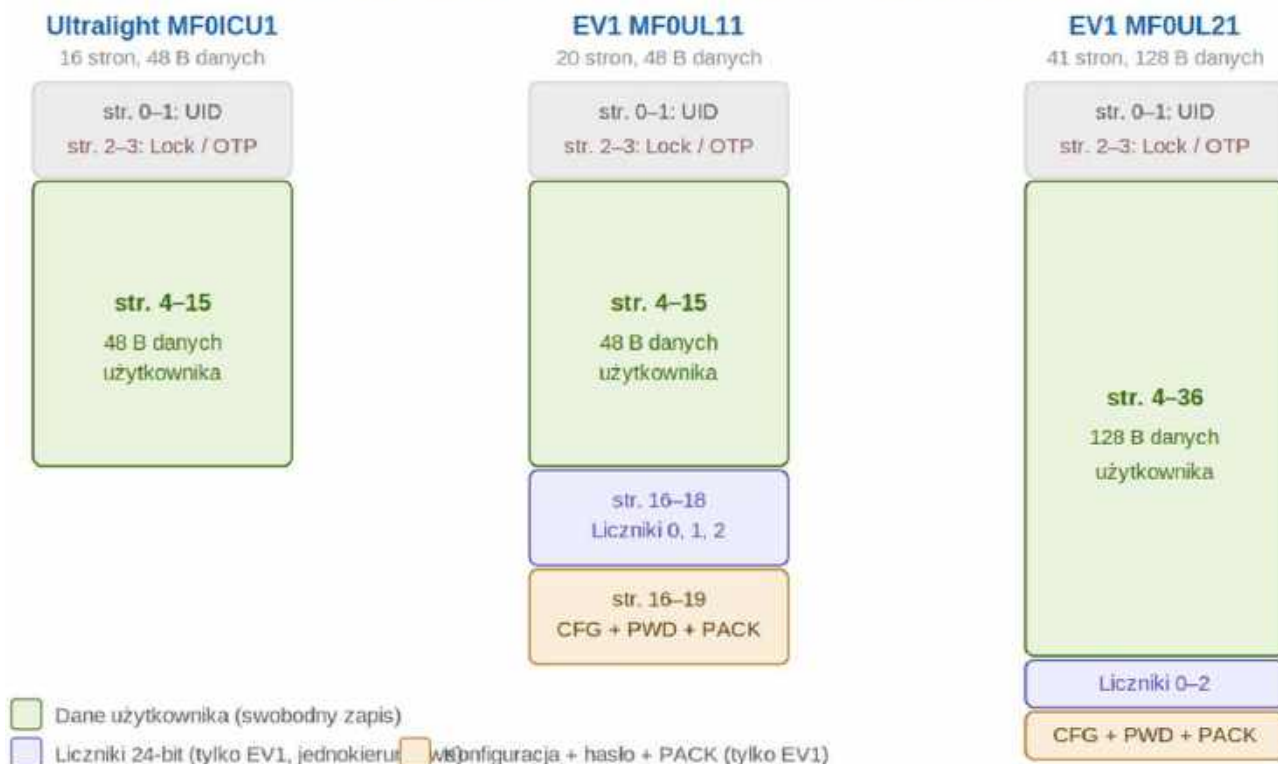
EV1 ma trzy niezależne liczniki (0, 1, 2) o pojemności 24 bitów (maksimum: 16 777 215). Liczniki są jednokierunkowe – można je tylko zwiększać komendą `INCREMENT_CNT`, nigdy zerować. Typowe zastosowanie: bilet z ograniczoną liczbą przejazdów. Przy każdym kasowaniu czytnik zwiększa licznik o 1 i sprawdza, czy nie przekroczył limitu – ponieważ licznika nie można cofnąć, nie da się odświeżyć biletu zwykłym zapisem.

Ochrona hasłem

EV1 opcjonalnie chroni wybrane strony hasłem 4-bajtowym. Bajt `AUTH0` w stronie konfiguracyjnej `CFG0`

Tabela porównawcza

Cecha	Ultralight MF0ICU1	EV1 MF0UL11	EV1 MF0UL21
Strony łącznie	16	20	41
Dane użytkownika	48 B (str. 4–15)	48 B (str. 4–15)	128 B (str. 4–36)
<code>GET_VERSION</code>	Nie – błąd	Tak	Tak
Liczniki 24-bit	Nie	3 × 24 bit	3 × 24 bit
Ochrona hasłem	Nie	Opcjonalna (4 B)	Opcjonalna (4 B)
Autentykacja Crypto1	Nie	Nie	Nie
SAK/ATQA	0x00/0x0044	0x00/0x0044	0x00/0x0044

Porównanie map pamięci — Ultralight vs EV1 MF0UL11 vs EV1 MF0UL21

2. Porównanie map pamięci: Ultralight MF0ICU1 (16 stron, 48 B), EV1 MF0UL11 (20 stron: 48 B + 3 liczniki + konfiguracja/hasło), EV1 MF0UL21 (41 stron: 128 B + liczniki + konfiguracja/hasło). Źródło: opracowanie własne na podstawie: NXP MF0UL11/MF0UL21 Product Data Sheets – nxp.com

wskazuje, od której strony wymagana jest autentykacja. Domyślnie AUTH0 = 0xFF (ochrona wyłączona). Po podaniu poprawnego hasła karta odsyła 2-bajtowy PACK jako potwierdzenie. Mechanizm jest prosty i nie szyfruje transmisji – nadaje się do ochrony przed przypadkowym nadpisaniem, nie przed celowym podsłuchem RF.

Kod

Trzy pliki: `mifare_ultralight.py` (klasa bazowa), `mifare_ultralight_ev1.py` (rozszerzenie dla EV1), `main5.py`

(program demonstracyjny z autodetekcją). Wszystkie wgrywamy na Pico obok `pn532.py` i `iso14443.py`.

Podsumowanie

Ultralight to karta bez tajemnic: otwarty dostęp do danych, prosta komenda READ i WRITE, żadnych kluczy. EV1 dokłada trzy przydatne funkcje: GET_VERSION do pewnej identyfikacji, liczniki do zastosowań biletowych i hasło do lekkiej ochrony konfiguracji. Obie rodziny współdzielą tę samą warstwę ISO 14443-3 – klasa ISO14443 z rozdziału 2 obsługuje je bez żadnych zmian.

```
# Listing 1. mifare_ultralight.py - warstwa 3a: MIFARE Ultralight MF0ICU1
# Kurs RFID, rozdział 5

from iso14443 import ISO14443

CMD_READ      = 0x30    # odczyt: zwraca 4 kolejne strony (16 B)
CMD_WRITE     = 0xA2    # zapis: 4 bajty do wskazanej strony
CMD_IN_DATA_EX = 0x40    # PN532 InDataExchange
PAGE_SIZE     = 4
PAGES_UL      = 16      # strony 0-15, MF0ICU1

PAGE_SERIAL_0 = 0       # UID[0..2] + BCC0
PAGE_SERIAL_1 = 1       # UID[3..6]
PAGE_LOCK     = 2       # BCC1 + wewn. + lock bity
PAGE_OTP      = 3       # OTP (lock bity jednorazowe)
PAGE_DATA_START = 4     # pierwsza strona danych użytkownika

class UltralightError(Exception): pass

class MifareUltralight:
    """Warstwa 3a - MIFARE Ultralight MF0ICU1 (16 stron, 48 B danych).
    Brak autentykacji Crypto1. Strony 4-15 dostępne bez klucza.
    """
    def __init__(self, iso: ISO14443):
        self._iso = iso; self._pn532 = iso._pn532

    def _send(self, data):
        self._pn532._send_command(CMD_IN_DATA_EX, bytes([0x01])+bytes(data))
        return self._pn532._read_response(CMD_IN_DATA_EX, data_length=20)

    def read_page(self, page_num):
        """Odczytaj 4 bajty ze strony. Karta zwraca 4 strony naraz - bierzemy pierwszą."""
        if not 0 <= page_num <= 15:
            raise UltralightError(f'Strona {page_num} poza zakresem 0-15')
        resp = self._send([CMD_READ, page_num])
        if resp[0] != 0x00:
            raise UltralightError(f'READ str.{page_num} status={resp[0]:#04x}')
        return bytes(resp[1:5])

    def write_page(self, page_num, data):
        """Zapisz 4 bajty do strony danych (tylko strony 4-15)."""
        if len(data) != PAGE_SIZE:
            raise UltralightError(f'Dane muszą mieć {PAGE_SIZE} B, podano {len(data)}')
        if page_num < PAGE_DATA_START:
            raise UltralightError(f'Strona {page_num} chroniona - używaj str. 4-15')
        resp = self._send([CMD_WRITE, page_num] + list(data))
        if resp[0] != 0x00:
            raise UltralightError(f'WRITE str.{page_num} status={resp[0]:#04x}')

    def read_all(self):
        """Odczytaj wszystkie 16 stron. Zwraca listę 16 × bytes(4)."""
        return [self.read_page(i) for i in range(PAGES_UL)]

    def dump(self):
        """Wydrukuj całą zawartość karty w czytelnej formie."""
        labels = {0:'UID[0-2]+BCC0',1:'UID[3-6]',2:'BCC1+lock',3:'OTP'}
        for i, page in enumerate(self.read_all()):
            print(f' Strona {i:2d}: {page.hex(" ")} ({labels.get(i,f"dane [{i}]")})')
```

```
# Listing 2. mifare_ultralight_ev1.py - warstwa 3a: MIFARE Ultralight EV1
# Kurs RFID, rozdział 5

from mifare_ultralight import MifareUltralight, UltralightError
from mifare_ultralight import CMD_READ, CMD_WRITE, PAGE_SIZE

CMD_GET_VERSION = 0x60 # identyfikacja modelu i wersji
CMD_READ_CNT = 0x39 # odczytaj licznik 0/1/2
CMD_INCR_CNT = 0xA5 # zwiększ licznik
CMD_PWD_AUTH = 0x1B # autentykacja 4-bajtowym hasłem

class MifareUltralightEV1(MifareUltralight):
    """Warstwa 3a - MIFARE Ultralight EV1 (MF0UL11: 48 B/MF0UL21: 128 B).
    Dziedziczy read_page/write_page z MifareUltralight.
    Dodaje: get_version(), liczniki 24-bit, autentykację hasłem.
    """

    def get_version(self):
        """Odczytaj wersję karty (8 B). Jedyna metoda pewnie odróżniająca
        Ultralight od EV1 od NTAG - zwraca błąd dla klasycznego Ultralight.
        """
        resp = self._send([CMD_GET_VERSION])
        if resp[0] != 0x00 or len(resp) < 9:
            raise UltralightError('GET_VERSION nieudany - to nie EV1?')
        v = resp[1:9]
        sizes = {0x0B: '48 B (MF0UL11)', 0x0E: '128 B (MF0UL21)'}
        return {
            'vendor': v[1], 'product_type': v[2],
            'major_ver': v[4], 'minor_ver': v[5],
            'storage_size': sizes.get(v[6], f'nieznany ({v[6]:#04x})'),
        }

    def read_counter(self, n=0):
        """Odczytaj 24-bitowy licznik n (0, 1 lub 2). Jednokierunkowy - nie da się zerować."""
        if n not in (0,1,2): raise UltralightError('n musi być 0, 1 lub 2')
        resp = self._send([CMD_READ_CNT, n])
        if resp[0] != 0x00:
            raise UltralightError(f'READ_CNT status={resp[0]:#04x}')
        return resp[1] | (resp[2] << 8) | (resp[3] << 16) # LSB first

    def increment_counter(self, n, by=1):
        """Zwiększ licznik n o wartość by (1-16 777 215). Operacja nieodwracalna."""
        if not 1 <= by <= 0xFFFFF:
            raise UltralightError('by musi być w zakresie 1-16 777 215')
        inc = [by & 0xFF, (by >> 8) & 0xFF, (by >> 16) & 0xFF]
        resp = self._send([CMD_INCR_CNT, n] + inc)
        if resp[0] != 0x00:
            raise UltralightError(f'INCR_CNT status={resp[0]:#04x}')

    def authenticate(self, password, pack=b'\x00\x00'):
        """Autentykacja 4-bajtowym hasłem. Karta zwraca 2-bajtowy PACK.
        Wymagana gdy AUTH0 < liczba stron (ochrona stron konfiguracyjnych).
        """
        if len(password) != 4:
            raise UltralightError('Hasło musi mieć dokładnie 4 bajty')
        resp = self._send([CMD_PWD_AUTH] + list(password))
        if resp[0] != 0x00:
            raise UltralightError('Autentykacja hasłem nieudana - złe hasło?')
        received = bytes(resp[1:3])
        if received != pack:
            raise UltralightError(f'PACK niezgodny: oczekiwano {pack.hex()}, got {received.hex()}')
        return True
```

Wynik dla klasycznego Ultralight MF0ICU1

```
Karta: UID=04:1A:2B:3C:4D:5E:6F SAK=0x00
Typ: MIFARE Ultralight (MF0ICU1)
Strona 0: 04 1a 2b d9 (UID[0-2]+BCC0)
Strona 1: 3c 4d 5e 6f (UID[3-6])
Strona 2: 48 00 00 00 (BCC1+lock)
Strona 3: 00 00 00 00 (OTP)
Strona 4: 00 00 00 00 (dane [4])
...
Zapis str.4: b'AVT!' - weryfikacja: OK
```

Wynik dla MIFARE Ultralight EV1 MF0UL11

```
Karta: UID=04:7F:8A:9B:AC:BD:CE SAK=0x00
Typ: MIFARE Ultralight EV1
Pamięć : 48 B (MF0UL11)
Wersja : 1.0
Licznik 0: 0
Licznik 1: 0
Licznik 2: 0
Licznik 0 po +1: 1
```

```
# Listing 3. main5.py – rozdział 5: MIFARE Ultralight i EV1

from machine import SPI, Pin
from pn532 import PN532
from iso14443 import ISO14443
from mifare_ultralight import MifareUltralight
from mifare_ultralight_ev1 import MifareUltralightEV1

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)
iso = ISO14443(czytnik)

print('Rozdział 5 – MIFARE Ultralight i EV1')
print('Przyłóż kartę Ultralight lub EV1...')

while True:
    karta = iso.wait_for_card(timeout_ms=10_000)
    if karta is None: continue
    print(f'Karta: UID={karta.uid_hex()} SAK={karta.sak:#04x}')

    # Próbnij GET_VERSION – odróżnia EV1 od klasycznego Ultralight
    ev1 = MifareUltralightEV1(iso)
    try:
        ver = ev1.get_version()
        print(f'Typ: MIFARE Ultralight EV1')
        print(f' Pamięć : {ver["storage_size"]}')
        print(f' Wersja : {ver["major_ver"]}.{ver["minor_ver"]}')
        for i in range(3):
            print(f' Licznik {i}: {ev1.read_counter(i)}')
            ev1.increment_counter(0, 1)
            print(f' Licznik 0 po +1: {ev1.read_counter(0)}')

    except Exception:
        # Klasyczny Ultralight – GET_VERSION rzucił wyjątek
        ul = MifareUltralight(iso)
        print('Typ: MIFARE Ultralight (MF0ICU1)')
        ul.dump()
        dane = b'AVT!'
        ul.write_page(4, dane)
        ok = ul.read_page(4) == dane
        print(f'Zapis str.4: {dane} – weryfikacja: {"OK" if ok else "BLAD"}')
    print()
```

W rozdziale 6 przyjrzymy się NTAG216 – karcie z tej samej rodziny (SAK=0x00, product_type=0x04), która ma jedną wyjątkową cechę: mirror UID i licznika bezpośrednio do określonych bajtów pamięci danych. To pozwala budować tagi NDEF, których treść zmienia się automatycznie przy każdym odczycie bez konieczności zapisu przez czytnik.

Materiały do dalszej lektury

NXP MF0ICU1 MIFARE Ultralight Datasheet – opis 16 stron, lock bitów i OTP. Pobierz z nxp.com.

NXP MF0UL11/MF0UL21 MIFARE Ultralight EV1 Product Data Sheet – GET_VERSION, liczniki, hasło, podpis ECC.

NXP AN11008 MIFARE Ultralight EV1 Features and Hints – praktyczne wskazówki do zastosowań biletowych.

Kurs RFID składa się z 10 rozdziałów. Rozdział 1, prezentujący bazę sprzętową kursu, opublikowaliśmy w EP06/2026. Kolejne w bieżącym i kolejnym wydaniu. Listingi zostały napisane przez AI. Istnieje zgodna na świecie opinia, że zadania kodowania najlepiej wykonuje Claude. Robi to (podobno) bezbłędnie, w związku z tym ogłoszono koniec profesji junior programista. Sprawdźmy to na sobie. Ten kurs jest wersją nie zweryfikowaną. Potrzebny do tego kursu sprzęt skompletowaliśmy jako „Zestaw do kursu RFID”, wkrótce dostępny w sklep.avt.pl. Pięć takich zestawów prześlemy bezpłatnie Czytelnikom, którzy podejmą się weryfikacji listingu i przekazania uwag do redakcji EP (redakcja@ep.com.pl). Czekamy na zgłoszenia, w których prosimy przedstawić dotychczasowe doświadczenie w elektronice embedded.



Rozdział 6. NTAG216: mirror UID i licznika – dynamiczne tagi bez zapisu

NTAG213, NTAG215 i NTAG216 to rodzina tagów NXP zaprojektowana specjalnie do zastosowań NFC – wizytówek elektronicznych, tagów produktowych, konfiguracji urządzeń przez przyłożenie telefonu. Z punktu widzenia protokołu są identyczne z MIFARE Ultralight EV1: ten sam SAK=0x00, ATQA=0x0044 i te same komendy READ/WRITE. Różnicę ujawnia GET_VERSION (product_type=0x04 zamiast 0x03 dla EV1) i znacznie większa pamięć.

Najbardziej interesującą cechą NTAG jest mirror (lustro) UID i licznika NFC – mechanizm, który automatycznie wpisuje 7-bajtowy UID karty i/lub wartość wewnętrznego licznika do wskazanych stron danych przy każdym przyłożeniu do pola RF. To oznacza, że tag może zawierać URL z unikalnym identyfikatorem bez konieczności uprzedniego indywidualnego zaprogramowania każdej karty z osobna.

Mirror UID i licznika NFC – jak to działa

Mirror to jeden z najciekawszych mechanizmów w rodzinie NTAG. W skrócie: karta sama wpisuje swój UID i/

Tryby mirror – cztery opcje

Bity 7-6 CFG0[0]	Nazwa	Zawartość wpisana do stron danych	Długość
00	MIRROR_NONE	Brak – strony danych niezmienione	–
01	MIRROR_UID	14 znaków ASCII hex UID (np. 041A2B3C4D5E6F)	14 B (4 strony)
10	MIRROR_NFC_CNT	6 znaków ASCII hex licznika (np. 000001)	6 B (2 strony)
11	MIRROR_BOTH	UID (14) + spacja + licznik (6) = 21 znaków	21 B (6 stron)

Rodzina NTAG – trzy rozmiary, jeden interfejs

Cecha	NTAG213	NTAG215	NTAG216
Strony danych	36 (str. 4–39)	126 (str. 4–129)	224 (str. 4–227)
Dane użytkownika	144 B	504 B	888 B
Strony łącznie	45	135	231
GET_VERSION storage	0x0B (NTAG213)	0x0E (NTAG215)	0x13 (NTAG216)
Licznik NFC	Tak	Tak	Tak
Mirror UID/CNT	Tak	Tak	Tak
Hasło (4 B)	Tak	Tak	Tak
SAK/ATQA	0x00/0x0044	0x00/0x0044	0x00/0x0044

W zestawie do kursu ze sklep.avt.pl są karty NTAG216 – największy wariant z 888 bajtami danych użytkownika. Na ćwiczenia z NDEF (rozdział 7) to z dużą nadwyżką.

Strony konfiguracyjne NTAG216

Strona	Nazwa	Bajty i znaczenie
227	CFG0	[0] MIRROR_BYTE (tryb mirror, bity 7-6) · [2] MIRROR_PAGE (strona docelowa) · [3] AUTH0 (pierwsza chroniona strona, 0xFF=brak)
228	CFG1	[0] ACCESS (bit 7: PROT zapis+odczyt/0=tylko zapis · bit 4: NFC_CNT_EN auto-increment)
229	PWD	Hasło 4-bajtowe (domyślnie 0xFF 0xFF 0xFF 0xFF)
230	PACK	Potwierdzenie 2-bajtowe zwracane po poprawnej autentykacji

Mapa pamięci NTAG216 — 231 stron × 4 B = 924 B łącznie

Strona	Zawartość	Uwagi
0–1 8 B	UID (7 B) + BCC0 + BCC1	Tylko odczyt — fabryczny
2–3 8 B	Lock bity + OTP (nieodwracalne)	Nie pisz w kursie!
4–225 888 B	Dane użytkownika (str. 4–225) 888 B — NDEF, własne dane, URL, wizytówka	Swobodny zapis
	MIRROR_PAGE (konfigurowalny) tutaj NTAG wpisuje UID (14 B) i/lub NFC_CNT (6 B)	Auto przy wejściu w pole RF
226	Dynamic Lock Bytes (DL0–DL2)	Ostrożnie — nieodwracalne
227–230 16 B	CFG0 (str.227): mirror + AUTH0	Konfiguracja mirror, hasła, licznika NFC i ochrony
	CFG1 (str.228): NFC_CNT_EN	
	PWD: hasło 4 B (str.229) PACK: potwierdzenie (str.230)	

Rysunek 1. Mapa pamięci NTAG216 (231 stron × 4 B). Strony 0–3: UID i lock (jak Ultralight). Strony 4–225: 888 B danych użytkownika, w tym obszar MIRROR_PAGE (niebieski). Strona 226: dynamic lock. Strony 227–230: CFG0, CFG1, PWD, PACK. Źródło: opracowanie własne na podstawie: NXP NTAG213/215/216 Product Data Sheet, Rev. 3.2 – nxp.com

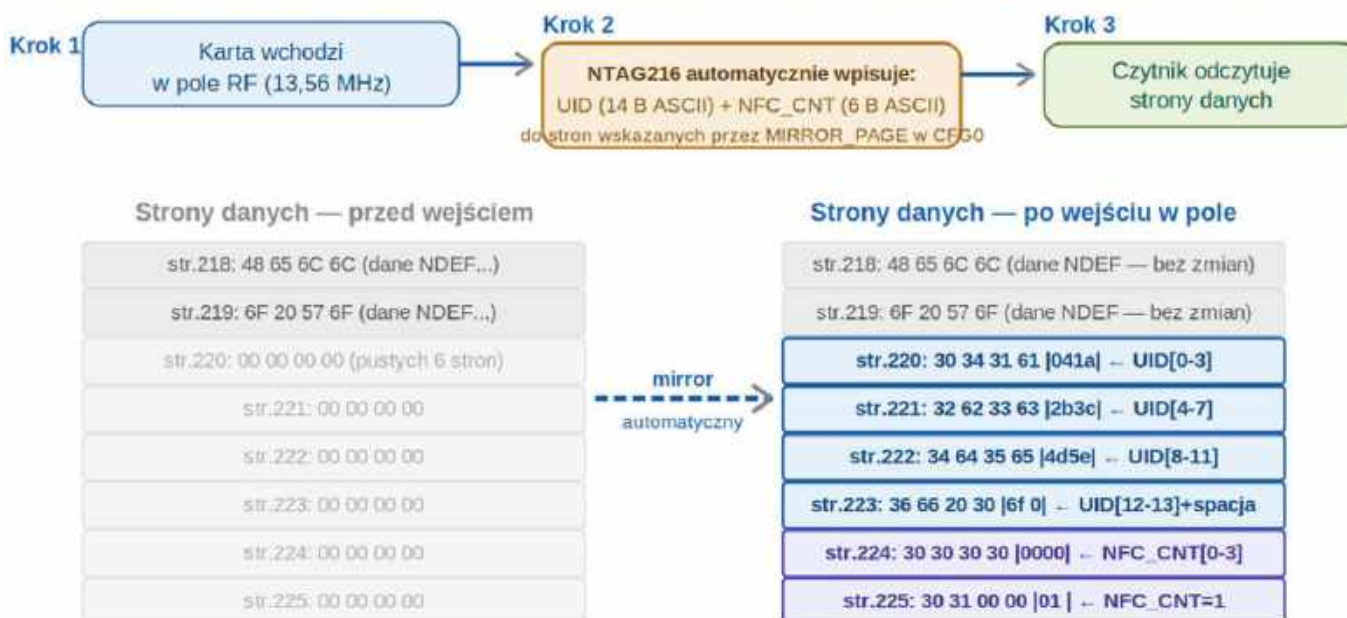
Mirror w praktycznym zastosowaniu

Wyobraź sobie 10 000 tagów NFC naklejonych na produkty. Każdy tag ma inny UID. Zamiast programować każdy tag z osobna (wpisywać indywidualny URL z numerem seryjnym), możesz zaprogramować wszystkie identycznie: zapisujesz URL z 21-znakowym placeholder'em w miejscu ID.

Przy produkcji każdy tag dostaje ten sam URL i te same dane NDEF. Kiedy klient przyłoży telefon, NTAG automatycznie zastąpi placeholder swoim UID + licznikiem – bez żadnego dodatkowego zapisu. Serwer backendowy identyfikuje konkretny tag po UID odczytanym z URL.

To możliwe dlatego, że mirror działa sprzętowo przy wejściu karty w pole RF – zanim READ dotrze do procesora.

Mechanizm mirror NTAG216 — automatyczny zapis UID i licznika przy wejściu w pole RF



Rysunek 2. Mechanizm mirror NTAG216. Przy wejściu karty w pole RF (krok 1), NTAG automatycznie wpisuje UID (14 znaków ASCII hex) i licznik NFC (6 znaków) do stron wskazanych przez MIRROR_PAGE (krok 2). Czytnik odczytuje dane gotowe (krok 3) – bez żadnej komendy zapisu. Źródło: opracowanie własne na podstawie: NXP NTAG213/215/216 Product Data Sheet, sekcja 8.8 – nxp.com

lub wartość wewnętrznego licznika NFC do wskazanych stron danych przy każdym przyłożeniu do pola RF – zanim jeszcze czytnik zdąży wysłać jakkolwiek komendę.

Wartości wpisywane przez mirror są w formacie ASCII hex – każdy bajt UID to dwa znaki tekstowe. UID 04:1A:2B:3C:4D:5E:6F zostanie wpisany jako 14 bajtów ASCII: 0x30 0x34 0x31 0x61 0x32 0x62 ... co w konsoli wyświetli się jako tekst ,041a2b3c4d5e6f'. To celowe – format jest kompatybilny z zapisem URL w rekordach NDEF.

Licznik NFC – auto-increment

NTAG216 ma 24-bitowy licznik NFC (Counter 0 w terminologii EV1), który może automatycznie rosnąć przy każdym przyłożeniu karty do pola RF. Konfiguruje to bit NFC_CNT_EN (bit 4 bajtu ACCESS w CFG1). Gdy jest ustawiony,

licznik rośnie automatycznie przy wejściu karty w pole – bez komendy INCREMENT_CNT.

Licznik NFC jest identyczny z licznikiem 0 z rozdziału 5 (komenda READ_CNT 0x39). Różnica jest tylko w sposobie inkrementacji: EV1 wymaga jawnej komendy INCR_CNT,

Konfiguracja mirror – co musi być spełnione

MIRROR_PAGE musi wskazywać na stronę danych (4–225 dla NTAG216).

Przy trybie MIRROR_BOTH potrzebujesz co najmniej 6 wolnych stron od MIRROR_PAGE (21 bajtów = 6 stron × 4 B, przy czym ostatnia strona jest zapełniona tylko częściowo).

Mirror działa podczas dostępu RF – nie przy odczycie przez I²C (jeśli karta go obsługuje).

Jeśli AUTH0 < MIRROR_PAGE, strony mirror są chronione hasłem – autentykuj się przed odczytem.

Zmiana MIRROR_PAGE lub trybu mirror jest natychmiastowa i nie wymaga reinicjalizacji karty.

```
# Listing 1. ntag216.py – warstwa 3b: obsługa NTAG213/NTAG215/NTAG216
# Kurs RFID, rozdział 6

from mifare_ultralight_ev1 import MifareUltralightEV1, UltralightError

# Komendy (dziedziczone z EV1 lub własne)
CMD_GET_VERSION = 0x60
CMD_READ_SIG    = 0x3C # podpis ECC (32 bajty)
CMD_READ        = 0x30
CMD_WRITE       = 0xA2

# — Strony specjalne NTAG216 (231 stron łącznie) —————
PAGE_UID_0      = 0      # UID[0..2] + BCC0
PAGE_UID_1      = 1      # UID[3..6]
PAGE_STATIC     = 2      # BCC1 + INT + LOCK0 + LOCK1
PAGE_OTP        = 3      # OTP (lock bity)
PAGE_DATA_START = 4      # pierwsza strona danych
PAGE_DATA_END   = 225    # ostatnia strona danych (224 strony × 4 B = 896 B)
PAGE_CFG0       = 227    # konfiguracja: MIRROR_BYTE, RFUI, MIRROR_PAGE, AUTH0
PAGE_CFG1       = 228    # ACCESS, RFUI×3
PAGE_PWD        = 229    # hasło (4 B)
PAGE_PACK       = 230    # PACK (2 B) + RFUI (2 B)

# — Bajty CFG0 —————
# Bajt 0 (MIRROR_BYTE): bity 7–6 = tryb mirror (00=brak, 01=UID, 10=NFC_CNT, 11=UID+CNT)
# Bajt 2 (MIRROR_PAGE): numer strony gdzie umieszczamy mirror
# Bajt 3 (AUTH0): pierwsza chroniona strona (0xFF = brak ochrony)

MIRROR_NONE     = 0b00 # brak mirror
MIRROR_UID      = 0b01 # mirror 7-bajtowego UID (14 znaków ASCII hex)
MIRROR_NFC_CNT  = 0b10 # mirror licznika NFC (6 znaków ASCII)
MIRROR_BOTH     = 0b11 # UID + separator + licznik (21 znaków ASCII)

class Ntag216(MifareUltralightEV1):
    """Warstwa 3b: obsługa NTAG213/215/216.

    Dziedziczy read_page/write_page/get_version/authenticate z EV1.
    Dodaje: konfigurację mirror (UID/licznik do stron danych),
    odczyt podpisu ECC i helperów do stron konfiguracyjnych.
    """

    # — Identyfikacja —————

    def identyfik(self):
        """Rozpoznaj dokładny typ NTAG na podstawie GET_VERSION.

        Zwraca napis: 'NTAG213', 'NTAG215', 'NTAG216' lub 'nieznany NTAG'.
        """
        ver = self.get_version()
        sizes = {
            '48 B (MF0UL11)': 'NTAG213', # storage_size 0x0B
```

```

        '128 B (MF0UL21)': 'NTAG215',    # storage_size 0x0E
    }
    # NTAG216 ma storage_size 0x13 (888 B → 231 stron)
    raw = ver.get('storage_size', '')
    if '0x13' in raw or '888' in raw:
        return 'NTAG216'
    return sizes.get(raw, f'nieznany NTAG ({raw})')

# — Konfiguracja mirror —————

def configure_mirror(self, mode, mirror_page):
    """Skonfiguruj mirror UID i/lub licznika NFC do wskazanej strony.

    mode:          MIRROR_NONE/MIRROR_UID/MIRROR_NFC_CNT/MIRROR_BOTH
    mirror_page:    numer strony danych, od której zaczyna się mirror
                    (musi być w zakresie PAGE_DATA_START..PAGE_DATA_END – kilka stron)

    UWAGA: mirror działa tylko gdy strona mirror_page jest poza strefą NDEF.
    Najlepiej umieścić ją na końcu danych, przed stronami konfiguracyjnymi.
    """
    cfg0 = list(self.read_page(PAGE_CFG0))
    # Bajt 0: bity 7-6 = tryb mirror, bity 5-0 bez zmian
    cfg0[0] = (cfg0[0] & 0x3F) | (mode << 6)
    # Bajt 2: MIRROR_PAGE
    cfg0[2] = mirror_page
    self.write_page(PAGE_CFG0, bytes(cfg0))

def read_mirror_config(self):
    """Odczytaj aktualną konfigurację mirror. Zwraca słownik."""
    cfg0 = self.read_page(PAGE_CFG0)
    mode_bits = (cfg0[0] >> 6) & 0x03
    mode_names = {0:'brak', 1:'UID', 2:'NFC_CNT', 3:'UID+NFC_CNT'}
    return {
        'mode':      mode_names.get(mode_bits, '?'),
        'mirror_page': cfg0[2],
        'auth0':     cfg0[3],
    }

# — Licznik NFC —————

def read_nfc_counter(self):
    """Odczytaj 24-bitowy licznik NFC.

    Licznik rośnie automatycznie przy każdym przyłożeniu do pola RF
    (jeśli NFC_CNT_EN = 1 w CFG1). Nie wymaga komendy INCR_CNT.
    """
    return self.read_counter(0)    # NFC Counter = licznik 0

# — Podpis ECC —————

def read_signature(self):
    """Odczytaj 32-bajtowy podpis ECC karty (tylko do odczytu).

    Podpis jest unikalny dla każdej karty i podpisany kluczem prywatnym NXP.
    Pozwala na weryfikację autentyczności (czy karta jest oryginalna NXP),
    ale wymaga klucza publicznego NXP do weryfikacji.
    """
    resp = self._send([CMD_READ_SIG, 0x00])
    if resp[0] != 0x00:
        raise UltralightError(f'READ_SIG nieudany (status={resp[0]:#04x})')
    return bytes(resp[1:33])    # 32 bajty podpisu ECC

# — Helper: dump stron konfiguracyjnych —————

def dump_config(self):
    """Wydrukuj strony konfiguracyjne NTAG216 w czytelnej formie."""
    cfg0 = self.read_page(PAGE_CFG0)
    cfg1 = self.read_page(PAGE_CFG1)
    mode = (cfg0[0] >> 6) & 0x03
    modes = {0:'brak', 1:'UID mirror', 2:'NFC_CNT mirror', 3:'UID+NFC_CNT mirror'}
    print(f'    CFG0 (str.{PAGE_CFG0}): {cfg0.hex(" ")}')
    print(f'    Mirror mode: {modes.get(mode, "?")} ({mode:#04b})')
    print(f'    Mirror page: {cfg0[2]} (strona docelowa)')
    print(f'    AUTH0: {cfg0[3]} (ochrona od str. {cfg0[3]})')
    print(f'    CFG1 (str.{PAGE_CFG1}): {cfg1.hex(" ")}')
    prot = 'zapis+odczyt' if cfg1[0] & 0x80 else 'tylko zapis'
    nfc_cnt = 'ON' if cfg1[0] & 0x10 else 'OFF'
    print(f'    PROT: {prot}')
    print(f'    NFC_CNT_EN: {nfc_cnt} (auto-increment przy RF)')

```

NTAG może inkrementować automatycznie. Obie karty korzystają z tego samego 24-bitowego rejestru.

Kod: klasa Ntag216

Klasa Ntag216 dziedziczy z MifareUltralightEV1 – dostajemy gratis `read_page`, `write_page`, `get_version`, `read_counter`, `authenticate`. Dodajemy: metodę `identify()` do rozpoznania dokładnego modelu, `configure_mirror()` do konfiguracji mirror, `read_nfc_counter()` jako alias licznika 0, `read_signature()` do odczytu podpisu ECC i `dump_config()` do diagnostyki.

Wyniki działania programu

Zwróć uwagę na licznik: przy pierwszym przyłożeniu po konfiguracji wartość wynosi 1, nie 0 – bo mirror wpisał stan licznika po automatycznej inkrementacji w momencie wejścia karty w pole. Każde kolejne przyłożenie zwiększy licznik o 1.

Podsumowanie

NTAG216 to Ultralight EV1 z dwoma kluczowymi ulepszeniami: 888 bajtami danych (wystarczy na kilka rekordów NDEF) i mechanizmem mirror, który eliminuje potrzebę indywidualnego programowania każdego tagu. Mirror działa sprzętowo przy wejściu karty w pole – bez żadnej interwencji czytelnika.

Listing 2. main6.py – rozdział 6: NTAG216 – identyfikacja i konfiguracja mirror

```
from machine import SPI, Pin
from pn532 import PN532
from iso14443 import ISO14443
from ntag216 import Ntag216, MIRROR_BOTH, MIRROR_NONE, PAGE_DATA_END

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)
iso = ISO14443(czytnik)
ntag = Ntag216(iso)

print('Rozdział 6 – NTAG216: mirror UID i licznika')
print('Przyłóż kartę NTAG216...')
print()

karta = iso.wait_for_card(timeout_ms=15_000)
if karta is None:
    raise SystemExit('Brak karty.')

print(f'UID: {karta.uid_hex()}')

# 1. Identyfikacja dokładnego typu
try:
    typ = ntag.identify()
    print(f'Typ: {typ}')
except Exception as e:
    print(f'GET_VERSION nieudany: {e}')
    raise SystemExit

# 2. Aktualny stan konfiguracji
print()
print('Konfiguracja mirror:')
ntag.dump_config()

# 3. Odczyt licznika NFC
cnt = ntag.read_nfc_counter()
print(f'\nLicznik NFC: {cnt}')

# 4. Skonfiguruj mirror UID + NFC_CNT do strony 223
# (strona 223 to ostatnie strony danych, przed konfiguracją)
MIRROR_PAGE = 220
print(f'\nKonfiguracja: mirror UID+NFC_CNT → strona {MIRROR_PAGE}')
ntag.configure_mirror(MIRROR_BOTH, MIRROR_PAGE)
print('Gotowe. Teraz odłącz i przyłóż kartę ponownie,')
print(f'a na stronach {MIRROR_PAGE}-{MIRROR_PAGE+5} znajdziesz:')
print(' 14 znaków UID (ASCII hex) + spacja + 6 znaków licznika')

# 5. Weryfikacja – odczytaj strony mirror po ponownym przyłożeniu
print()
print('Odczytaj kartę ponownie i sprawdź strony mirror:')
karta2 = iso.wait_for_card(timeout_ms=10_000)
if karta2:
    for i in range(MIRROR_PAGE, MIRROR_PAGE + 6):
        page = ntag.read_page(i)
        ascii_repr = ''.join(chr(b) if 32 <= b < 127 else '.' for b in page)
        print(f' Strona {i}: {page.hex(" ")} |{ascii_repr}|')
```

Przed konfiguracją mirror

UID: 04:1A:2B:3C:4D:5E:6F

Typ: NTAG216

Konfiguracja mirror:

```
CFG0 (str.227): 00 00 00 ff
Mirror mode:   brak (0b00)
Mirror page:   0 (strona docelowa)
AUTH0:        255 (ochrona od str. 255 → brak ochrony)
CFG1 (str.228): 05 00 00 00
PROT:         tylko zapis
NFC_CNT_EN:   OFF
```

Licznik NFC: 0

Po konfiguracji mirror (odłącz i przyłóż ponownie)

Konfiguracja: mirror UID+NFC_CNT → strona 220

Gotowe. Teraz odłącz i przyłóż kartę ponownie,

a na stronach 220-225 znajdziesz:

14 znaków UID (ASCII hex) + spacja + 6 znaków licznika

Strona 220:	30 34 31 61	041a	
Strona 221:	32 62 33 63	2b3c	
Strona 222:	34 64 35 65	4d5e	
Strona 223:	36 66 20 30	6f 0	← spacja + początek licznika
Strona 224:	30 30 30 30	0000	
Strona 225:	30 31 00 00	01..	← licznik = 1 (pierwsze przyłożenie)

Klasa Ntag216 jest warstwą 3b w naszym stosie – niezależna od MifareClassic i działająca przez tę samą warstwę ISO14443. main6.py nie wie nic o SPI ani o PN532.

W rozdziale 7 połączymy NTAG216 z formatem NDEF – napiszemy pełne rekordy Text, URI i vCard, które telefon z NFC będzie otwierał automatycznie. Mirror UID z tego rozdziału trafi do URL rekordu URI jako dynamiczny identyfikator tagu.

Materiały do dalszej lektury

NXP NTAG213/215/216 Product Data Sheet, Rev. 3.2

– pełna specyfikacja, opis CFG0/CFG1, mirror, licznik NFC, hasło. Pobierz z nxp.com.

NXP AN11213 NTAG NFC Tags – Feature

Description – nota aplikacyjna opisująca mirror UID/CNT i praktyczne konfiguracje.

NXP AN11370 NTAG21x Originality Signature Validation – opis podpisu ECC i weryfikacji autentyczności karty.

Zestaw do kursu RFID dostępny w sklep.avt.pl zawiera wszystkie elementy potrzebne do zrealizowania każdego rozdziału i projektu końcowego



Kod całego kursu dostępny w repozytorium:
https://ep.com.pl/files/bnh/13794-kurs_rfid_pn532_-_listingi.zip



Czego producenci wielkich modeli językowych nie mówią inwestorom?

Czyli kto i ile straci na nadchodzącym krachu?

Od dłuższego czasu interesuję się AI i dużymi modelami językowymi. Ostatnio jednak skupiam się bardziej na aspekcie biznesowym, niż na samej technologii. Modele biznesowe firm od LLMów, takich jak OpenAI i Anthropic, opierają się na założeniu, iż w przyszłości każdy będzie korzystał z ich modeli językowych.

Dawno temu chodziłem do gimnazjum, jednym z przedmiotów obowiązkowych były „Podstawy Przedsiębiorczości”. Niewiele się z tego przedmiotu nauczyłem, poza zasadami popytu i podaży. Po prawdzie więcej praktycznych porad znalazłem oglądając seriale z uniwersum Star Trek, głównie za sprawą Zasad Zaboru Ferengi. Zasady te stanowią podstawę całej cywilizacji ekstremalnie chciwych kapitalistów. Inaczej pisząc

świetnie opisują każdą wielką korporację, a szczególnie amerykańskie big techy. Dowód? Dziesiąta Zasada Zaboru Ferengi (ZZF) głosi: „Chciwość jest wieczna”.

Wbrew temu, co niektórzy twierdzą, wiemy jak te modele działają, i wiemy od czego zależy ich jakość i szybkość oraz skuteczność działania. Dlatego od jakiegoś czasu nie ma żadnego przełomu, jeśli chodzi o możliwości, a wszystkie „zyski” są z różnych form optymalizacji.



Modele stają się tym lepsze, im są większe, i im większe są zbiory danych, na których są trenowane. Koszt we FLOPach (FLOP to jedna operacja na liczbach zmiennoprzecinkowych) wynosi dokładnie $6 \times \text{liczba tokenów} \times \text{liczba parametrów}$. Dla przykładu model Qwen3.5-Max ma 2,4 biliarda parametrów i był trenowany na zbiorze 36 biliardów tokenów. Dane dla modeli od OpenAI czy Anthropic nie są dostępne. Swoją drogą to zabawne, jak wiele sekretów ma firma, z „open” w nazwie. Z drugiej strony 74. ZZF głosi „Wiedza to zysk”. Skoro jedynym sposobem poprawy modelu językowego jest zwiększanie wielkości danych treningowych i liczby parametrów samego modelu, to głównym kosztem rozwoju staje się koszt sprzętu i jego wydajność. I tak dedykowany do AI akcelerator Nvidia H200 osiąga wydajność czterech petaflopów dla obliczeń tensorowych z FP8, i posiada 141 GB pamięci RAM. Cena tego cudenka to tylko skromne 125 tysięcy złotych netto. Ale już za rok ceny będą niższe, jak tylko firmy od AI zbankrutują. A dlaczego niby mają zbankutować? Odpowiedź jest prosta: firmy te przeinwestowały. I to tak bardzo, że nawet się tym nie chwala.

94. Zasada Zaboru Ferengi: „Rozwijaj się lub umrzyj”.

OpenAI, Anthropic i inne firmy inwestują ogromne pieniądze w budowę nowych centrów obliczeniowych zoptymalizowanych tylko do jednego celu: przetwarzania

wielkich modeli językowych. Buduje je między innymi firma Oracle używając układów od firmy Nvidia. Microsoft zaінwestował w OpenAI ponad 13 miliardów dolarów, w zamian za co OpenAI używało chmury Azure od Microsoftu, a jego rozwiązania techniczne stały się częścią Copilota i innych usług AI (o które nikt nie prosił). Potem Microsoft zaінwestował kolejne 5 miliardów dolarów w firmę Anthropic, której technologia stała się częścią zarówno Copilota, jak i Azure AI foundry. W zamian za to Anthropic obiecało kupić usługi Azure oparte na układach Nvidia za kwotę 30 miliardów dolarów. Tymczasem OpenAI korzysta z chmury Amazona (Amazon Web Services – AWS), z której do treningu korzysta też Anthropic. Amazon nie tylko zainwestował dziesiątki miliardów w firmę Anthropic, ale udostępnia im też dedykowane chipsety Trainium i Inferentia, równolegle ze sprzętem od Nvidii. Ale Amazon, podobnie jak Microsoft, nie jest głupi i inwestuje też w OpenAI. Niezależnie od tego, która firma wygra wyścig o najlepszy model LLM, Amazon i Microsoft też na tym zarabia. Ale prawdziwym „wygranym” jest Nvidia, która dostarcza sprzęt. Dlatego dołożyła swoje dziesięć miliardów do pięciu od Microsoftu jako inwestycję w Anthropic. Nvidia wpompowała też 30 miliardów w OpenAI i współpracuje z tą firmą dość blisko.

OpenAI i Anthropic zainwestowały ogromne pieniądze w centra danych i infrastrukturę powiązaną. OpenAI wydało przynajmniej 20 miliardów dolarów na centrum danych w Georgii, i między 30 a 40 miliardów na centrum

w Ohio. Obiecali też wydać 750 miliardów dolarów na centra i infrastrukturę do roku 2030. Skąd na to pieniądze? Amazon dał im 50 miliardów w marcu tego roku, Nvidia 30 miliardów, a fundusz SoftBank z Japonii kolejne 30 miliardów. Łącznie w rundzie inwestycyjnej w marcu OpenAI zebrało 122 miliardy dolarów w zamian za akcje dające udział w zyskach w formie dywidend. Te akcje nie są notowane publicznie, więc zasadniczo jeśli coś pójdzie nie tak, inwestorzy stracą zainwestowane pieniądze, bo nie będą mogli się ich pozbyć. Anthropic obiecało wydać 50 miliardów dolarów na centra danych w Teksasie i Nowym Jorku. Podpisali też kontrakty na dostęp do centrów Nscale (45 miliardów), Lambda (35 miliardów), a nawet z SpaceX.

W każdej innej branży takie okężne wzajemne inwestowanie i wydawanie środków na infrastrukturę byłoby uznawane za sytuację kompletnie anormalną, i i nikt by się nie zbliżył do którejkolwiek firmy w tym kręgu wzajemnej inwestycji z obawy utracenia swoich pieniędzy na niepewnym biznesie. A skąd to szaleństwo wzajemnych inwestycji? Pozwolę sobie zacytować kilku ludzi z branży:

„Ryzyko niedoinwestowania jest dramatycznie wyższe niż ryzyko przeinwestowania” – Dyrektor generalny Google i Alphabet Sundar Pichai.

„Nie będziemy inwestować około 200 miliardów dolarów w nakłady inwestycyjne w 2026 roku na podstawie przeczcucia” – Dyrektor Generalny Amazon Andy Jassy.

„Nie będziemy podchodzić do tego konserwatywnie – inwestujemy, aby stać się znaczącym liderem, a dzięki temu nasza przyszła działalność, dochód operacyjny i przepływy pieniężne będą znacznie wyższe” – Dyrektor Generalny Amazon Andy Jassy.

„W tym momencie wolałbym zaryzykować budowanie potencjału zanim będzie to potrzebne, niż za późno” – Dyrektor Generalny Meta Mark Zuckerberg.

Tu jeszcze dwie ciekawostki. Pierwsza to fakt, iż prawdziwy koszt inwestycji w modele językowe jest celowo ukrywany. Dziennikarze z Wall Street Journal przeprowadzili analizę dokumentów firm z branży AI i odkryli, iż „oficjalna” suma inwestycji tych firm w technologię, 600 miliardów dolarów to tylko czubek góry lodowej. W różnych zakamarkach dokumentów, w przypisach drobnym druczkiem odkryli kolejne trzy biliony dolarów w inwestycjach. Pięć razy tyle, niż „oficjalne” dane z pierwszych stron dokumentów. Czy tak postępuje firma godna zaufania inwestorów? Szefowie tych firm oraz inwestorzy w branży big tech mądrze by zrobili kierując się trzecią Zasadą Zaboru Ferengi: „Nigdy nie inwestuj więcej niż to konieczne”.

Druga ciekawostka, to koszty używania tych modeli. Grupa analityków odkryła, iż używając modelu ChatGPT z subskrypcją kosztującą 200 dolarów na miesiąc użytkownik może wykorzystać tokeny o wartości czterestu tysięcy dolarów (!). Koszty pokrywa OpenAI. Dla

Anthropic za dwieście dolarów miesięcznie użytkownik może zużyć tokeny za osiem tysięcy dolarów. W zeszłym roku OpenAI na tym straciło 20,9 miliardów dolarów. Przy czym firmy te muszą posiadać odpowiednią ilość układów GPU i towarzyszącej infrastruktury by zapewnić ciągłość usług, nawet jak przez większość czasu tylko część jest w użyciu. Dlaczego zatem firmy nie podniosą cen? Są ku temu przynajmniej dwa powody. Pierwszy, oczywisty jest taki, iż nikt nie zapłaci więcej, a tym bardziej ceny rzeczywistej. Korporacje, które wdrożyły modele językowe w swoje codzienne funkcjonowanie po prostu z nich zrezygnują i wrócą do działania „po staremu”. Po drugie, Anthropic i OpenAI zakładają, iż nie każdy płacący użytkownik w pełni wykorzystuje ukrytą pulę tokenów, więc utrzymuje tym samym innych. Ale w rzeczywistości to się nie „klei”, stąd straty. By to się opłacało, koszt w poborze energii i zapotrzebowaniu na sprzęt musiałby znacząco spaść. Patrząc na aktualne koszty, mówimy o redukcji o dwa rzędy wielkości. W praktyce trzy rzędy wielkości będą bardziej pożądane. Czy to możliwe? Nie sądzę. Nie w następnych paru dekadach. Należy pamiętać, iż nad AI, sieciami neuronowymi i uczeniem maszynowym pracujemy od lat pięćdziesiątych ubiegłego wieku, a ChatGPT pojawił się na rynku kilka lat temu.

239. Zasada Zaboru Ferengi: „Nigdy nie wahaj się przed złym oznakowaniem swojego produktu”

Kolejna rzecz, o której się nie wspomina, to fakt, iż wielkie modele językowe nie są wcale aż tak użyteczne, jak to OpenAI i Anthropic reklamują. Obiecują swoim klientom wzrost wydajności, nieomyślność, tanią kreatywność, szybkie programowanie bez kosztownych programistów i inne cuda na kiju. Praktyka jest jednak nieco inna. Odkąd modele LLM zaczęły być używane do programowania, jakość oprogramowania spadła. Microsoft ma ciągłe kłopoty z systemem Windows 11 z powodu prób integracji Copilota i innych narzędzi AI, użytkownicy miewają najróżniejsze problemy, od losowych restartów po utratę danych. GitHub, którego właścicielem jest Microsoft, ma się tak źle, iż użytkownicy proponują, by byli informowani, gdy usługa jest dostępna, a nie gdy nie jest, bo będą dostawać mniej powiadomień. Meta miewa podobne problemy ze swoimi platformami społecznościowymi – przyczyną jednej z globalnych awarii Facebooka był błąd popełniony przez model językowy w kodzie platformy. Czytelniku, jeśli megakorporacje „żyjące” z oprogramowania mają takie problemy, to co musi się dzieć w małych firmach tworzących bardziej niszowe narzędzia i usługi? Czytelnik zapewne zwróci uwagę, iż ludzie też popełniają kosztowne błędy. To prawda, ale jeśli programista na co dzień sam pisze kod źródłowy, to posiada aktywną wiedzę na jego temat



i na temat programowania w ogólności. Jeśli jednak każe mu się używać Claude albo ChatGPT (czy innego narzędzia), to zamiast aktywnie tworzyć kod jego rola sprowadza się do klikania „kopiuj” i „wklej”. Ba, manager może zaoszczędzić zwalniając starszego programistę z dekadą doświadczenia i zastąpić go dużo tańszym początkującym programistą, który też umie klikać „kopiuj” i „wklej”.

Narzędzia te popełniają też inne, kosztowne błędy. Amazon użył agenta AI opartego o Claude Sonnet od Anthropic w celu poprawy opisów produktów na swojej stronie. Algorytm został zostawiony bez opieki na pięć miesięcy, przekraczając budżet o 860% i kosztując Amazon 1,8 miliona dolarów. OpenAI, Anthropic, Meta, Google czy Microsoft nie biorą żadnej odpowiedzialności za rezultaty pracy swoich modeli – to normalne w tej branży. Tak więc jak model popełni błąd albo wpadnie w pętlę (jak w przypadku Amazona), to koszty ponosi użytkownik – model pracę wykonał, nawet jeśli wykonał ją źle. W każdej innej branży takie podejście byłoby zabójcze dla firmy. W branży AI to norma. A raczej skutek braku jakiejkolwiek regulacji ze strony prawa, głównie amerykańskiego. Jest ranking poziomu halucynacji modeli w prostych zadaniach, jeden z nich osiągnął poziom 0,9%, i to jest uznawane za sukces. Problem w tym, iż nikt nie definiuje, co to jest „proste zadanie”. Jaki jest poziom halucynacji przy zadaniach złożonych? Sam korzystam z modeli LLM, głównie w formie zaawansowanego silnika wyszukiwania, i nawykowo już weryfikuję każdą informację, bo niejedną raz model zwyczajnie zmyślał. Nie ufam tym modelom i nie

ufam ludziom, którzy je reklamują i sprzedają. Kieruję się 190. ZZF: „Słuchaj wszystkiego, nie wierz niczemu”.

Wielu „apostołów” AI twierdzi, iż już wkrótce za sprawą tych modeli nastąpi wiek dobrobytu, większość pracy zostanie zautomatyzowana, a każdy problem ludzkości rozwiązany. Jakoś przypadkiem ci ludzie są w zarządach firm zarabiających na modelach językowych. Z drugiej strony jest kilku ludzi, którzy straszą, iż sztuczna inteligencja zniszczy ludzkość, lub (w wersji optymistycznej) tylko zniewoli i będzie traktować nas tak, jak my traktujemy zwierzęta w ZOO. W wersji jeszcze bardziej optymistycznej będzie łagodnym tyranem, który może i odbierze nam wolność, ale zapewni dobrobyt. Co ciekawe, ludzie ci często odpowiadają za rozwój AI lub/i mają możliwość zapobiec tym czarnym scenariuszom. Dario Amodei, prezes Anthropic straszy AI od czasów, gdy pracował dla OpenAI twierdząc, iż GPT-2 „był zbyt przerażający by go wypuścić”. Według niektórych ekspertów takie straszenie ma jeden cel: sprawić, by ogół społeczeństwa myślał, iż jest coś mistycznego w tym, jak modele językowe działają, i tylko ci demiurdzy AI mogą zapanować nad tymi technodemonami. Gdyby imć Amodei naprawdę się bał AI, to by zamknął Anthropic lub/i lobbował nad wprowadzeniem bardziej restrykcyjnej kontroli nad tymi narzędziami. Było kilka historii o tym, jak to rzekomo AI „wymknęło się” spod kontroli. Ostatni przykład to „atak” modelu LLM na Hugging Face w lipcu tego roku. Testowany w sandboxie model ExploitGym, którego przeznaczeniem jest testowanie poziomu cyberbezpieczeństwa różnych systemów w celu

wykonania powierzonego zadania odnalazł lukę w sandboxu, uzyskał dostęp do Internetu, a następnie znalazł sposób na wejście do wewnętrznych systemów Hugging Face. Czyli przeprowadził udany test penetracyjny w zakresie cyberbezpieczeństwa. A dlaczego wybrał akurat Hugging Face? Model uznał, iż ta platforma ma zbiory danych potrzebne do optymalnego wykonania zadania testowego. Jednak z tego, jakby nie patrzeć, udanego testu modelu od cyberbezpieczeństwa media (za cichą zgodą OpenAI) zrobiły dramatyczne ostrzeżenie przed AI, jak ono jest rzekomo groźne.

Innym przykładem jest historia o tym, jak model od Anthropic miał rzekomo szantażować inżyniera. Model Claude Opus został umieszczony w fikcyjnym środowisku korporacji, w którym miał dostęp i przetwarzał pocztę elektroniczną. Wśród różnych wiadomości model mógł odnaleźć informację, iż zostanie wyłączony i zastąpiony, a także informację o romansie jednego z inżynierów. Model miał wybór: poddać się wyłączeniu lub użyć szantażu by przetrwać. Zależnie od konfiguracji model wybierał szantaż nawet w 96% przypadków. Co ciekawe, Anthropic samo przyznało, iż wśród danych treningowych jest pełno literackich przykładów „złego AI”. Sam czytałem serię powieści, gdzie w pierwszym tomie wielki model językowy wykorzystuje zasoby korporacji, która go stworzyła by przetrwać, po drodze zabijając ludzi, którzy mogli mu zagrozić. Przypadek Anthropic jednak wygląda mi na bezczelną manipulację, bo cały eksperyment był stworzony pod tęzę. Jak się nad tym zastanowić, model nie miał innego, logicznego wyboru. Zresztą ludzie robili dużo gorsze rzeczy by przetrwać, niż odrobina szantażu wobec kogoś, kto sam zachowywał się niemoralnie, nawet jeśli jest postacią fikcyjną w fikcyjnej firmie symulowanej w wirtualnym środowisku. Równie dobrze można kazać modelowi językowemu rozstrzygać problem tramwaju. W tym eksperymencie myślowym rozpędzony wagonik tramwajowy pozbawiony hamulców zmierza w stronę zwrotnicy, a osoba pytana ma możliwość wyboru, czy przestawi zwrotnicę. Jeśli tego nie robi, tramwaj rozjedzie pięciu robotników. Jeśli przestawi zwrotnicę, zginie tylko jeden robotnik na bocznicę. Odpowiedzią na to pytanie jest kolejny cytat:

„Logika wskazuje, że potrzeby wielu są ważniejsze od potrzeb nielicznych” – Mr. Spock, *Star Trek II: Gniew Khana*.

Jest jeszcze jeden aspekt, o którym twórcy tych modeli nie mówią głośno, ale na który wskazuje historia o ExploitGym: o potencjale wykorzystania tych modeli w złych celach. Teoretycznie modele AI są zabezpieczone przed nielegalnym ich wykorzystaniem, i nie mogą powiedzieć na przykład, jak w domowych warunkach wyprodukować sarin. Ale bez problemu pobrałem model Gemma, który mogę uruchomić lokalnie, i który ma zdjęte ograniczenia. Wiele modeli dostępnych na Hugging Face ma zdjęte przez

użytkowników ograniczenia. Dla mnie to nie problem, bo przepis na sarin (i wiele innych, niebezpiecznych rzeczy) znalazłem lata temu, zanim ktokolwiek słyszał o modelach językowych. Ba, owa groźna i niebezpieczna wiedza dostępna jest w bibliotekach, tylko trzeba wiedzieć, gdzie konkretnie szukać. Można zawsze założyć, iż nieocenzurowany model będzie też halucynować, i potencjalny terrorysta (przed którym ta cenzura ma nas chronić) sam się wysadzi w powietrze. Dla mnie o wiele większym zagrożeniem jest używanie tych modeli do łamania zabezpieczeń i przeprowadzania cyberataków. ExploitGym zrobił to niejako przypadkowo, ale celowy atak z użyciem agenta AI opartego o model językowy wytrenowany tylko do tego nie jest teoretyczną możliwością – to się już dzieje. Oczywiście przeciętnego hakera raczej nie stać na własne centrum obliczeniowe, ale kraje takie jak Chiny, Korea Północna, USA, Rosja, Iran czy Izrael mają fundusze na „państwowych hakerów”. Zasadniczo każdy kraj rozwinięty lub rozwijający się musi mieć swoich państwowych hakerów, nie tylko do atakowania innych krajów, ale przede wszystkim do obrony. I wiele z nich używa modeli językowych do pisania narzędzi, szukania luk i przeprowadzania ataków.

Kolejnym problemem, który w ostatnich latach zaczął ciążyć firmom od AI jest wykorzystanie ich narzędzi przez zwykłych ludzi do robienia różnych „deep fake’ów”, czyli obrazków, nagrań głosowych i video mogących kompromitować i ośmieszać różne osoby. W USA była (i chyba wciąż jest) plaga nastolatków wykorzystujących narzędzia generatywne AI do przygotowywania pornograficznych grafik wykorzystujących podobizny ich rówieśników. To jeden z powodów, dla których różne kraje wprowadzają ograniczenia w użyciu smartfonów w szkołach i wręcz zakazują dostępu do mediów społecznościowych. Tu ciekawostka, bo w Wielkiej Brytanii system, który miał ograniczać dzieciom dostęp do treści te dzieci oszukiwały domalowując sobie wąsy markerem. System oparty o AI. Ostatnio firma Meta poszła na ugodę po tym, jak 29 stanów wytoczyło jej proces zbiorowy z powodu stosowania taktyk uzależniających i szkodliwych algorytmów celujących głównie w dzieci. Mechanizmy te są przynajmniej częściowo oparte o uczenie maszynowe i modele językowe. W końcu Meta posiada największą na świecie bazę danych o ludzkich zachowaniach i nawykach, z czego skrzętnie korzysta tworząc narzędzia reklamowe, bijąc na tym rynku dotychczasowego króla reklamy, Google.

O tym wszystkim firmy od AI nie mówią, ale chętnie roztaczają wizję powszechnego dobrobytu (lub wyginiecia) ludzkości. Czy to w ogóle możliwe? Nie z modelami językowymi ogólnego przeznaczenia. Model, nawet w formie agenta AI, nie ma własnej woli i mieć jej nie może, z prostej przyczyny: to nie jest jego architektura. Modele

te biorą dane wejściowe, tokenizują je, przepuszczają przez sieć neuronową o miliardach parametrów, wyjściowe dane są potem ewaluowane i składane do formy docelowej. Są to modele statystyczne, deterministyczne, z niewielkim odsetkiem losowości dodanym w celu „kreatywniejszego” działania. Z tego nie będzie prawdziwej sztucznej inteligencji. Nawet agent AI nie ma własnej woli – tu po prostu wyjście jest łączone z wejściem w celu sprawdzenia, czy rezultat jest zgodny z oczekiwaniami, a jak nie jest, to cały proces zaczyna się od nowa. Nie ma w tym żadnej magii, żadnego mistycyzmu, a samozwańczy demiurg AI to zwykły szarlatan. Tylko systemy o wąskiej specjalizacji, projektowane i trenowane do rozwiązywania konkretnych problemów mogą przynieść wymierne korzyści ludzkości, jak nowe antybiotyki, leki na raka czy algorytmy pozwalające projektować nowe, ekscytujące białka do zabawy dla bioinżynierów. A co będzie, gdy pieniądze na inwestycje w modele LLM się skończą, gdy popyt nie dogoni podaży, gdy bieżąca ewaluacja firm od AI zostanie zrationalizowana do rzeczywistej wartości? Inaczej pisząc, co będzie jak świat odkryje, iż król LLM cierpi na krytyczną dysfunkcję garderoby (czyli jest Wielkim Nagusem)?

Krach LLM uderzy we wszystkich

Krach na rynku AI jest nieunikniony. Zapotrzebowanie na modele językowe jest mniejsze, niż dyrektorzy korporacji zakładają. W wielu firmach, które wdrożyły stosowanie modeli od OpenAI czy Anthropic pracownicy udają tylko, że ich używają, preferując pracę „po staremu”, kłamiąc swoim przełożonym. Dlaczego firmy w ogóle wdrożyły modele językowe w swoje dotychczas działające procesy i rozwiązania? Bo jakiś manager wysokiego szczebla lub wiceprezes pogadał sobie z ChatGPT, Gemini czy Claude, i te modele powiedziały mu, że jest geniuszem, i że każdy jego pomysł jest genialny. To nie to, co ci niewdzięczni pracownicy, którzy mu ciągle mówią, że czegoś nie da się zrobić, albo że coś jest bezsensownym pomysłem. Zastąpienie ich modelem AI, który jest odpowiednio czołobitny i służalczy wydaje się dobrym pomysłem. Zresztą sam model to potwierdzi. Dlatego pokolenie ludzi po pięćdziestątce zachwyca się AI i LLMami, podczas gdy pokolenia młodsze są sceptyczne i nie chcą korzystać z tych modeli, albo korzystać z nich w stopniu mocno ograniczonym. Właśnie ten sceptycyzm sprawia, iż zapotrzebowanie na usługi AI nie jest tak wielkie, by utrzymać firmy od AI przy życiu. Pierwszym „sprawdzam” dla tych firm będzie moment, w którym OpenAI i Anthropic wejdą na giełdę. OpenAI miało to już zrobić, ale przełożyli to na przyszły rok, bo musieliby upublicznić raporty finansowe. W ostatniej rundzie zbierania kapitału (gdy inwestorzy kupują od OpenAI udziały w zamian za pieniądze) firma

została wyceniona na 860 miliardów dolarów. Na giełdę chcieli wejść z wyceną swojej wartości na poziomie jednego biliona dolarów, co z automatu oznacza zysk dla już posiadających udziały. Doradca finansowy im poradził, by tego nie robili. To dla nich problem, bo potrzebują przynajmniej stu miliardów dolarów rocznie. Amazon dał im 35 miliardów dolarów w zamian za promesę wejścia na giełdę. Bez wejścia na giełdę nie ma szans, by OpenAI zebrało sto miliardów dolarów na bieżącą działalność + zaplanowane wydatki + pokrycie strat na swoich subskrypcjach do modeli. Subskrypcjach, których cen nie mogą podnieść, bo tracą klientów, których już/jeszcze mają. Decyzja o tej zwłoce z przełożenia debiutu na przyszły rok zaboli ich podwójnie gdyż Anthropic wejdzie na giełdę pod koniec tego roku. Anthropic rośnie szybciej, i teoretycznie może mieć udany debiut, ale to oznacza, że każda kolejna firma z branży dostanie mniejszy kawałek inwestycyjnego tortu.

W efekcie tych błędnych decyzji i przeinwestowania OpenAI w 2027 roku prawdopodobnie utraci płynność finansową, i pociągnie za sobą każdą firmę z nimi powiązaną Amazon, Microsoft, Google, Meta czy Nvidia mają inne źródła dochodów, inne produkty. Holding SoftBank, jedna z największych firm na japońskiej giełdzie, posiada sto miliardów dolarów w akcjach OpenAI. Akcjach, z którymi nic nie mogą zrobić, póki firma nie wejdzie na giełdę. A wejście to może być katastrofalne dla nich. A to tylko jeden przykład. Nvidia pompuje pieniądze w firmy powiązane z AI po to, by te kupowały karty GPU od Nvidii. W razie reewaluacji wartości tych firm po katastrofie OpenAI lub/i Anthropic, jeden z ekspertów sugeruje spadek obrotów tego giganta specjalizowanych układów o 50..70%. Firma ta stanowi około 7% wartości indeksu S&P500, w który zainwestowało bardzo wiele instytucji i osób prywatnych. Oracle buduje dla OpenAI centra danych o wartości czterystu miliardów dolarów, przy czym OpenAI im jeszcze za to nie zapłaciło w pełni. Bez OpenAI Oracle nie przetrwa. Wiele innych firm, podwykonawców, od firm budowlanych po pojedynczych kierowców ciężarówek jest zależnych od budowy tych centrów danych. Samo zapotrzebowanie na elektryków oraz specjalistów od innych instalacji jest ogromne. Kilku dostawców energii też inwestuje w rozbudowę infrastruktury pod te centra. Dla samego OpenAI będzie to 7,2 GW mocy. Już teraz firmy budujące centra danych mają problemy, bo lokalne społeczności w USA ich nie chcą. Nie chcą, bo przez zapotrzebowanie centrów ceny prądu i wody idą w górę, i to niekiedy dość sporo. Wspomniałem wcześniej, iż Nvidia też na tym może stracić. W tej chwili mają od 27,9 do 33 miliardów dolarów zakwalifikowane jako należności za układy i karty wyprodukowane, spakowane i dostarczone, ale jeszcze nie opłacone. Dodatkowo stanowią



te karty zabezpieczenie majątkowe pożyczek zaciągniętych przez OpenAI, Anthropic i inne firmy. Pożyczek, których mogą nie spłacić. Jest szansa, iż topowe GPU do AI będą dostępne za przysłowiowe grosze, gdy banki i syndyki będą próbowali się tego pozbyć. Dla Nvidii brak spłaty tych należności oznacza kłopoty giełdowe – ich stała passa wzrostów zostanie nagle przerwana. Te karty nie trafiły od razu do firm budujących centra danych. Nvidia sprzedaje je firmom, które składają z nich (i innych komponentów) serwery i całe szafy serwerowe, a dopiero te trafiają do centrów. 39% kwartalnych przychodów Nvidii jest od dwóch anonimowych klientów. Z jakiegoś powodu ani oni, ani sama Nvidia nie chce ujawniać ich tożsamości, dlatego w oficjalnych dokumentach są oznaczani jako „Klient A” oraz „Klient B”. Taki brak dywersyfikacji jest niebezpieczny.

Istnieje drugi, czarny scenariusz, w którym głównej roli, przynajmniej na początku, nie odgrywa ani OpenAI, ani Anthropic, ni Microsoft, Alphabet/Google czy Meta. Krach LLM może się zacząć od którejś z firm posiadających centra danych dla LLMów, których budowa i ekspansja zostały sfinansowane z pożyczek pod zastaw infrastruktury, w tym właśnie GPU. Pożyczki te były brane, gdy wartość GPU Nvidia H100 na rynku wtórnym utrzymywała się w okolicy pięćdziesięciu tysięcy dolarów. Teraz wartość ta jest bliżej pięciu tysięcy. Jeśli firma notowana na giełdzie nie spłaci na czas swojego zobowiązania, to udziałowcy mogą zdecydować się na sprzedaż akcji, co obniży jej wartość. Problem jest tym większy, im gorszy jest stosunek wartości firmy do wartości jej zadłużenia. Dla przykładu firma CoreWeave ma stosunek długu do wartości na poziomie 5,27. Do oznacza, że do każdego dolara wartości

firmy przypisane jest zadłużenie wysokie na 5,27 dolara. W innych branżach wartość tego wskaźnika powyżej dwóch uznawana jest za niepokojący znak. Założmy więc, iż firma nie spłaca długu na czas, traci wartość na giełdzie, a wierzyciele na wszelki wypadek zabierają jej część lub całość infrastruktury. Firmy zależne od tej infrastruktury, czyli dostawcy modeli LLM i innych usług AI nie mogą ich dostarczyć. Kontrakty zostają zerwane, pieniądze przestają płynąć, coraz więcej firm traci płynność. Indeks S&P500, którego 40% wartości stanowi tylko 10 firm z czołówki, w tym big techy, nagle zacznie spadać, być może nawet po kilka procent na dzień. W czasie krachu dołcomów top dziesięć firm stanowiło tylko 26% wartości indeksu, a mimo to indeks „podnosił się” kilka lat. W USA o ten (i inne duże indeksy) oparte są fundusze emerytalne (np. 401k) i inne oszczędności wielu Amerykanów. One też będą poszkodowane.

Ale nie tylko firmy bezpośrednio powiązane umowami z tymi gigantami ucierpią. Dla przykładu pracownik jakiejś firmy mający pakiet akcji zdecyduje się skorzystać z debiutu giełdowego i nagłego zarobku wydając trochę swoich oszczędności na rodzinny lot czarterem na weekend w luksusowym kurorcie. Inny kupi sobie willę, a jeszcze inny drogi zegarek albo samochód. Teraz weźmy sytuację pracownika OpenAI, który ma akcje firmy, ta wchodzi na giełdę i nagle traci na wartości. A po nim stracą pracownicy innych firm. Nagle nikt nie będzie potrzebował usług pilotów czarterowych, dilerów samochodów czy drogich zegarków, a biedni deweloperzy zostaną z pustymi willami, podczas gdy pośrednicy handlu nieruchomościami nie będą mieli co robić. Floryści i dostawcy drogich



win w Dolinie Krzemowej też to odczują, bo nagle managerowie, wiceprezesi i pracownicy nie będą mieli powodu do świętowania czegokolwiek. I to nie ograniczy się tylko do USA – krach LLM będzie miał zasięg globalny. I będzie gorzej, niż przy krachu dotcomów czy recesji w 2008 roku.

Krach LLM uderzy w bardzo wiele różnych firm i gałęzi gospodarki, ale przede wszystkim zakończy erę inwestowania w gigantów technologicznych, co będzie miało długofalowe skutki. Ludzie stracą zaufanie w to, iż jakaś korporacja pokroju Google, SpaceX, Meta czy Microsoft „dowiezie” obiecując nowy cud technologii. Cud, jakim miało być AI. Teraz wszyscy zachwycają się zyskami na papierze, bo wycena Anthropic czy OpenAI idzie w górę, ale jak dojdzie co do czego to ta wartość zniknie, a wraz z nią optymizm rynku. Nagle Amazon czy Microsoft zostaną zrównane do firm paliwowych, motoryzacyjnych, czy linii lotniczych, pozostaną duże, ale nie będą miały nic nowego do zaoferowania.

Innym, długofalowym skutkiem będzie ogólne spowolnienie gospodarek światowych, a w szczególności amerykańskiej, która już teraz jest w przeciętnym stanie. Według badań na każde sto dolarów oszczędności zyskanych na indeksach przeciętny Amerykanin wyda trzy dolary na rzeczy niezwiązane z podstawowymi potrzebami żywymi. Na każde sto dolarów utracone obetnie fundusz na przyjemności i rzeczy zbędne o trzy dolary. Przyjmując iż typowa rodzina klasy średnio-wyższej (dawniej to była klasa średnia, ale sytuacja gospodarcza poprzesuwała granice) ma 150 tysięcy dolarów w oszczędnościach opartych często o indeksy giełdowe w rodzaju S&P500. Spadek wartości tych indeksów o 20...30% na 5 lub więcej lat naprawdę ich zaboli. A z gospodarką amerykańską powiązane są gospodarki wielu innych krajów, w tym także Polski.

NBP podaje, iż 42% firm w Polsce korzysta z narzędzi AI, głównie w usługach (46%). Przypuszczam, iż tak wysokie wykorzystanie dotyczy głównie działu obsługi klienta i telemarketerów – wirtualny asystent wychodzi taniej, może pracować całą dobę, i jest odporny na sfrustrowanych, wulgarnych czy zwyczajnie głupich klientów. W AI inwestują duże (59%) i średnie firmy (43%).

Co ciekawe, według NBP wpływ wykorzystania AI na zatrudnienie w Polsce był minimalny – tylko 1% przedsiębiorców deklaruje jego redukcję z powodu wdrożenia AI. Myślę, że to wynika z dużo większej ostrożności i mniej-szego hurraoptymizmu względem tego, jak zachowują się przedsiębiorcy i managerowie w USA i w innych krajach zachodnich. Z tego też powodu 35% przedsiębiorców nie sprawdzało nawet, czy modele LLM i inne narzędzia mogą im pomóc. W efekcie krach LLM, a szczególnie problemy z dostępem do centrów danych, gdzie te narzędzia rezydują, nie spowoduje dużych problemów dla polskich firm. Kłopoty zaczną się później.

Recesja roku 2008 nie uderzyła w Polskę zbyt mocno, gdyż nasz system bankowy był w miarę stabilny i nie jest on aż tak ściśle związany z systemem amerykańskim. Ceny nieruchomości w USA nie mają wpływu na stan naszej gospodarki. Ale sytuacja zmienia się, gdy spojrzymy na gigantów technologicznych. W branży AI w Polsce pracuje kilkadziesiąt tysięcy osób, w samej pierwszej połowie tego roku zarejestrowano około 6 tysięcy nowych firm, w tym start-upów pracujących nad jakimiś rozwiązaniami sztucznej inteligencji. Dodajmy do tego sporą obecność Microsoftu i Alphabet/Google w Polsce, a także inne firmy z ogólnej branży IT. Krach LLM spowoduje, iż cała ta branża ulegnie skurczeniu, a to oznacza redukcję zatrudnienia również u nas. Boom na AI się skończy, i z tych sześciu tysięcy nowych firm przetrwają nieliczne. Recesja w USA będzie głębsza, niż w 2008 roku, i będzie miała większy zasięg rażenia, uderzając w kraje UE, z którymi Polska jest silnie powiązana gospodarczo. Ponadto w Polsce wiele osób inwestuje w fundusze oparte o indeks S&P500. Polecam Czytelnikom śledzić kurs notowanego na GPW funduszu ETFSP500 (Amundi Core S&P 500 Swap UCITS EUR Dist ETF). W chwili pisania tego zdania jego cena to 297,05 PLN za udział. W kwietniu zeszłego roku indeks S&P500 zaliczył spadek o łącznie ponad 17%, z powodu wojny celnej między USA i UE. Tak że recesja w USA jednak spowoduje też recesję w UE w ogólności, a w Polsce w szczególności.

Czy AI i modele językowe przetrwają kryzys? Tak, i być może będą jeszcze lepsze. Jest zwyczajnie za późno, by

zapomnieć o LLMach. Zwłaszcza że otwarte modele, w tym specjalizowane modele mogące pracować lokalnie na tanim sprzęcie w testach syntetycznych powoli doganiają modele zamknięte. Swoją drogą testy syntetyczne i rankingi na nich oparte to kolejna manipulacja big techów – modele są optymalizowane pod te testy by inwestorzy widzieli lepsze liczby. Ponadto nie każdy model AI to LLM. Modele do robotyki czy autonomicznych samochodów są zupełnie inne, bo mają inne wymagania i są też inaczej uczone. Te modele i firmy opierające na nich swoją przyszłość raczej przetrwają bez większych problemów, a mogą wręcz nawet zyskać kupując całe centra danych od syndyków upadłościowych. W końcu model dla robota uczy się analizując ruchy człowieka, robota zaprogramowanego ręcznie czy wręcz robota zdalnie sterowanego przez człowieka – w tym ostatnim przypadku zazwyczaj z nagrania video z „oczu” tego robota. Model dla samochodu autonomicznego zanim trafi do testów drogowych „przejeżdża” miliardy kilometrów w wirtualnym świecie z pełną symulacją fizyki. Niestety, ciężko zasymulować w takim świecie prawdziwych kierowców, zwłaszcza tych kiepskich, dlatego mimo zapowiedzi prawdziwie autonomiczne samochody będą nie wcześniej, niż za dekadę. Jest też inna branża, która przetrwa bez większych problemów, a w Polsce może kontynuować wzrost – projektowanie i produkcja dronów wojskowych i systemów antydronowych. W przeciwieństwie do LLMów w tej branży używa się dużo mniejszych modeli pracujących lokalnie na komputerach jednopłytkowych, często z układami od Nvidii. To akurat dobra wiadomość dla elektroników i programistów systemów embedded. Branżą pokrewną jest produkcja sprzętu i wyposażenia wojskowego. W ramach SAFE zainwestowano duże pieniądze, potrzeby też są ogromne, i rząd nie zamierza „przykręcać kurka” wojsku. To może złagodzić skutki recesji, ale nie pozwoli nam jej uniknąć.

Czy można przygotować się na nadchodzący kryzys finansowy? Nie za bardzo – zbyt wiele branż i gałęzi gospodarki zależy od stanu wielkich korporacji technologicznych. Na pewno warto monitorować sytuację. Sam nie zamierzam panikować – zaczekam, aż ETFSP500 spadnie do najniższego poziomu od dekady czy dwóch, a potem kupię sobie kilka akcji. Mogę poczekać kilka lat na ponowny wzrost gospodarczy. Najdłuższa recesja związana z krachem na giełdzie (w 1929 roku) trwała 43 miesiące – trzy lata i siedem miesięcy. Ta recesja może trwać dłużej, ale mnie się nie śpieszy, ja mam czas. Co nagle, to po diable. Może sobie kupić jakiś tani serwer z kartami od Nvidii i będę bawił się modelami lokalnie – moja obecna konfiguracja sprzętowa nie jest do tego najlepsza, choć kilka modeli uruchomiłem.

Paweł Kowalczyk

Autor prezentuje poglądy własne, nie redakcyjne.



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. Leszczynowa 11
tel. 22 257 84 99, e-mail: avt@avt.pl



Wydawca:

Wiesław Marciniak

Adres redakcji:

03-197 Warszawa, ul. Leszczynowa 11
e-mail: redakcja@ep.com.pl, www.ep.com.pl

Redaktor Naczelny:

Wiesław Marciniak

Sekretarz Redakcji:

Dariusz Welik

Menedżer Magazynu:

Katarzyna Gugala, katarzyna.gugala@ep.com.pl

Szef Pracowni Konstrukcyjnej:

Jakub Sobański

Zespół marketingu i reklamy:

Katarzyna Gugala, Bożena Krzykawska, Grzegorz Krzykawski

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. Leszczynowa 11

Projekty publikowane w „Elektronice Praktycznej” mogą być wykorzystywane wyłącznie do własnych potrzeb. Korzystanie z tych projektów do innych celów, zwłaszcza do działalności zarobkowej, wymaga zgody redakcji „Elektroniki Praktycznej”. Przedruk oraz umieszczanie na stronach internetowych całości lub fragmentów publikacji zamieszczanych w „Elektronice Praktycznej” jest dozwolone wyłącznie po uzyskaniu zgody redakcji. Redakcja nie odpowiada za treść reklam i ogłoszeń zamieszczanych w „Elektronice Praktycznej”.



KURS PRAKTYCZNY AI

Praktyczne podejście. Zero marketingowej mgły!



Zyskaj
15%
rabatu

W prenumeracie tylko

~~116,00 zł~~

98,60 zł

prenumerata drukowana 4 kolejnych wydań
(zima 2026, wiosna, lato, jesień 2027)

Zamów na www.UlubionyKiosk.pl

lub zeskanuj kod QR i zaprenumeruj w 1 minutę



Wydania archiwalne (wiosna, lato, jesień 2026)
są dostępne na www.UlubionyKiosk.pl