



ELEKTRONIKA

dla wszystkich

nr 02/2026 (361) • luty • www.elportal.pl

PROJEKTY dla elektroników

- ▶ PICO – miniaturowy analizator elektroakustyczny
- ▶ 262 144 sposoby na „Grę w życie”
- ▶ Generator sygnału testowego FM na pasmo 2 m

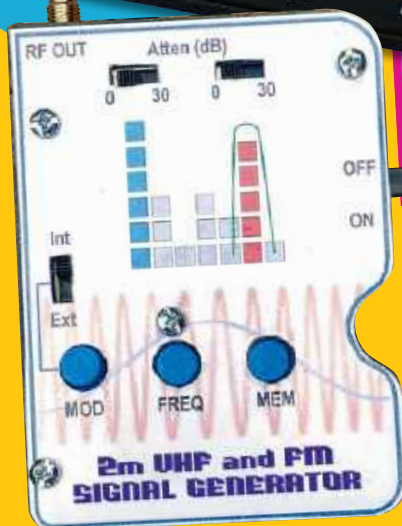
DIY dla wszystkich

- ▶ Prosty i niezawodny czujnik temperatury
- ▶ Generator funkcji na Raspberry Pi

TUTORIALE

- ▶ Bezprzewodowy Internet o globalnym zasięgu
- ▶ Chirurgia obwodowa: Zniekształcenia i obwody zniekształcające, część 3
- ▶ Ekscytacje Maxa: Migające diody LED i śliniący się inżynierowie, część 29
- ▶ Audio OUT: Uniwersalna płytką z pojedynczym wzmacniaczem operacyjnym, zoptymalizowana do zastosowań w układach akustycznych, część 1

PICO miniaturowy analizator elektroakustyczny



Generator
sygnału
testowego FM
na pasmo 2 m



Pomocna dłoń



automatykaB2B.pl

EP.com.pl

Największy
portal dla
elektroników
konstruktorów

eprasa.pl/fe6a48ff3f



FIRMA PIEKARZ
CZĘŚCI ELEKTRONICZNE

przełączniki
półprzewodniki
złącza
przekładniki
radiatory
obudowy
i wiele więcej...

www.piekarz.pl

Subscribe to Elektor's newsletter and get the chance to

WIN

a Raspberry Pi Pico W board



www.elektor.com/eda



Subscribe to Elektor's newsletter, get a €5 coupon code and get the chance to WIN a Raspberry Pi Pico W board



Be one of the 10 fortunate winners!



elektor
design > share > earn

-15%
NA START
192,80 zł

-30%
po pierwszym roku
prenumeraty
158,80 zł

-40%
po drugim roku
prenumeraty
136,10 zł

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

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

Rozpocznij swoją przygodę z *Elektroniką dla Wszystkich*. Decydując się teraz na roczną prenumeratę drukowaną, otrzymasz nie tylko dostęp do najnowszych wydań, ale i **znakomity start dzięki zniżce 15%** na pierwsze zamówienie!

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

Im dłużej jesteś z nami, tym więcej oszczędzasz:

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

Jak otrzymać rabat za lojalność?

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

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

Wybierz prenumeratę cyfrową PDF i ciesz się dostępem do czasopisma nawet 7 dni przed oficjalną premierą w kioskach. Oszczędzaj czas i pieniądze – skorzystaj z **rabatu 30%** na roczną e-prenumeratę w cenie 126,80 zł.

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

Zyskaj nieograniczony dostęp do zasobów dla pasjonatów elektroniki!

Tylko prenumeratorzy mają pełny dostęp do:

- cyfrowego archiwum *Elektroniki dla Wszystkich* na www.elportal.pl/archiwum
- projektów DIY+ na www.elportal.pl/diy

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

ARCHIWUM



Zacznij korzystać z pełnych zasobów już dziś!

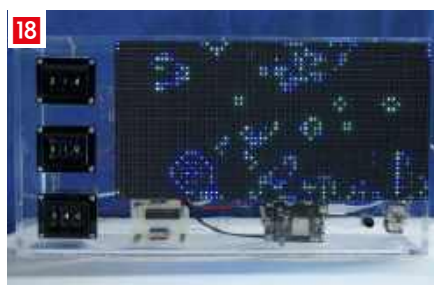


8

Projekty dla elektroników:

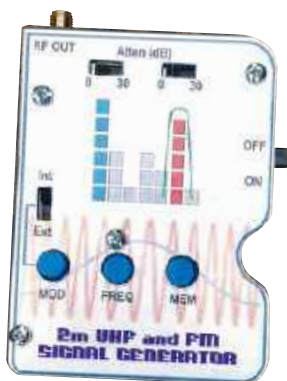
PICO – miniaturowy analizator elektroakustyczny	8
262 144 sposoby na „Grę w życie”	18
Generator sygnału testowego FM na pasmo 2 m.....	24

Tutoriale:



18

Bezprzewodowy Internet o globalnym zasięgu	34
Chirurgia obwodowa:	
Zniekształcenia i obwody zniekształcające, część 3	46
Ekscytacje Maxa:	
Migające diody LED i śliniacy się inżynierowie, część 29	52
Edukacja w EdW dla szkół i uczelni: Wykład 38 – Kwantyzacja.....	56



24

Audio OUT:

Uniwersalna płytką z pojedynczym wzmacniaczem operacyjnym, zoptymalizowana do zastosowań w układach akustycznych, część 1	63
--	----

DIY dla wszystkich:

Prosty i niezawodny czujnik temperatury.....	69
Generator funkcyjny na Raspberry Pi.....	72

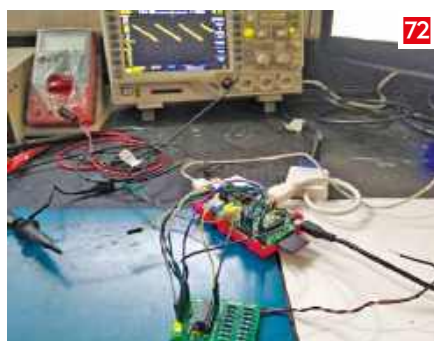


69

Junior:

Dwudzieste spotkanie z najmłodszymi pasjonatami elektroniki	77
---	----

Na zdjęciu na okładce Tymek – Młodzi Entuzjaści Elektroniki, Wrocław



72

DIY PLUS

tylko dla prenumeratorów zamawiających prenumeratę na www.UlubionyKiosk.pl

7-segmentowy mini zegar wykorzystujący PIC16F628A i DS1307 RTC	90
Światło LED oparte na czujniku zbliżeniowym	90

Rubryki stałe:

Prenumerata	3
Od redakcji.....	5
Poczt.....	7

A za miesiąc w marcowym EdW



Analogowy generator 1 kHz wysokiej jakości

Precyzyjny układ na wzmacniaczach operacyjnych OPA2211 i sprzężeniu optycznym, bez zbędnej elektroniki cyfrowej. Zapewnia wyjątkową czystość sygnału dzięki stabilizacji amplitudy i starannie dobranym komponentom. Całość zasilana jest napięciem $\pm 15\text{ V}$ i nie wymaga programowania. Ciekawa alternatywa dla drogich komercyjnych generatorów audio.

Analogowy regulator głośności sterowany cyfrowo

Nowoczesny układ, zaprojektowany jako bezpośredni zamiennik klasycznego potencjometru. Zapewnia precyzyjne ustawienie poziomu, doskonałe dopasowanie kanałów oraz wieloletnią niezawodność, eliminując problemy zużycia mechanicznego. Praktyczne rozwiązanie do wzmacniaczy i przedwzmacniaczy audio.

Krzemowy świerszcz

– e-grajek niczym ten z bajki

Znasz bajkę o koniku polnym, który całe lato nic tylko grał na swych skrzypcach? Niewielki układ oparty na mikrokontrolerze wiernie naśladuje cykanie świerszcza, wykorzystując przetwornik piezoelektryczny i sprytnie algorytmy czasowe. Reguluje na oświetlenie i losowo zmienia tempo dźwięków. Potrafi też naśladować żabę albo kanarka. Intrygujący gadżet do samodzielnego montażu.

Rewersyjny licznik częstotliwości

Sprytnie rozwiązanie problemu dokładnych pomiarów sygnałów wolnozmiennych. Za sprawą pomiaru okresu, zastosowanego w miejsce klasycznego zliczania impulsów, pozwala precyzyjnie mierzyć przebiegi zarówno bardzo wolne, jak i te nieco szybsze – o częstotliwości do około 10 MHz. Kwintesjencja nowoczesnej techniki pomiarowej.

Dla studentów i uczniów: garść technicznych tutoriali

Dla szukających inspiracji: ciekawe projekty DIY

Dla najmłodszych: kolejny zestaw z serii AVTEDU

**W kioskach
od 27 lutego**

Luty pod znakiem sygnałów

Luty kojarzy się z sercami, symbolami i deklaracjami. Z uczuciami, które – podobnie jak sygnały elektryczne – bywają ulotne, zniekształcone, czasem trudne do uchwycenia, a jednak bardzo realne. W elektronice nauczyliśmy się je mierzyć, analizować i filtrować. W życiu... bywa z tym różnie. Być może dlatego lutowy numer „Elektroniki dla Wszystkich” tak wyraźnie krąży wokół sygnałów – tych akustycznych, radiowych, cyfrowych i tych, które trudno opisać wzorem, ale łatwo poczuć.

Numer otwieramy projektem, który idealnie wpisuje się w ideę „narzędzia blisko ręki”. Miniaturowy analizator audio oparty na Raspberry Pi Pico to przykład nowoczesnej elektroniki w kieszonkowym wydaniu. Potrafi generować i analizować sygnały niskiej częstotliwości, pokazuje przebiegi, widmo i zniekształcenia, a dzięki własnemu wyświetlaczowi, złączom RCA i zasilaniu bateryjnemu może pracować zupełnie samodzielnie – także poza warsztatem. To urządzenie, które zachęca do eksperymentów i uczy, że pomiar nie musi oznaczać skomplikowanego stanowiska laboratoryjnego.

Z pomiarów audio płynnie przechodzimy w świat fal radiowych. Generator testowy CW/FM na pasmo amatorskie 2 m to projekt, który łączy sprytnie wykorzystanie syntezy DDS z praktycznym podejściem do testów odbiorników VHF. Aliasowanie, wąskopasmowe filtry LC, modulacja FM i obsługa CTCSS składają się tu na narzędzie niewielkie, ale bardzo konkretne w zastosowaniach. To dobry przykład inżynierskiej elegancji – maksimum funkcji przy minimum środków.

Nie samym pomiarem jednak elektronika żyje. Brian White zabiera Czytelnika w świat cyfrowej zabawy, która szybko przeradza się w wizualną fascynację. Gra w życie Conwaya w wersji sprzętowej, z możliwością wyboru każdej z 262 144 wersji zasad, pokazuje, jak proste reguły potrafią generować złożone, hipnotyzujące struktury. Kolorowa matryca LED i przemyślany algorytm sprawiają, że każdy eksperyment staje się małym dziełem sztuki – dowodem na to, że matematyka i elektronika potrafią być nie tylko użyteczne, ale i piękne.

Jak zawsze, w numerze nie zabrakło też inżynierskich dygresji i solidnej dawki doświadczenia. Max Maxfield z właściwym sobie humorem zagląda za kuliszy projektu animatronicznej głowy, pokazując, że prawdziwa nauka często zaczyna się tam, gdzie coś nie działa zgodnie z planem. Z kolei Jake Rothman wraca do fundamentów, prezentując uniwersalną płytkę z pojedynczym wzmacniaczem operacyjnym – prostą bazę do testów, eksperymentów i budowy klasycznych układów analogowych.

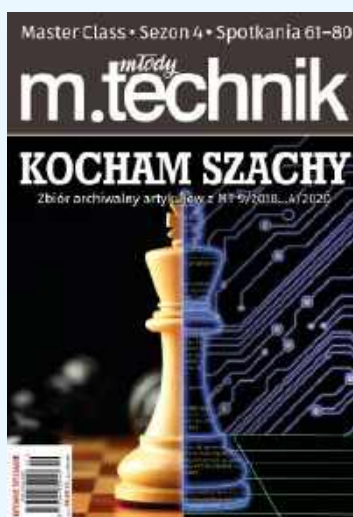
Dla tych, którzy lubią spojrzenie szersze i głębsze, przygotowaliśmy teksty prowadzące zarówno w przeszłość, jak i w przyszłość. David Maddison przybliży kuliszy systemów satelitarnych Starlink i Starshield, pokazując, jak nowoczesna technika zmienia globalną komunikację. Ian Bell w kolejnej odsłonie „Circuit Surgery” schodzi na poziom zniekształceń sygnału, ucząc, jak świadomie je kształtować – od analizy FFT po praktyczny odsłuch.

Szczegółne miejsce w lutowym numerze zajmuje dwudzieste spotkanie EdW Junior, tematycznie powiązane z walentynkami. Tym razem artykuł prowadzi Czytelnika znacznie dalej niż instrukcja montażu, a Bijące serce LED, reagujące na ciepło dłoni, staje się pretekstem do rozmowy o różnych obliczach miłości – również tej rodzinnej, opiekuńczej, przyjaźnielskiej i tej skierowanej ku własnym pasjom. Pokazuje, że nauka elektroniki może iść w parze z rozmową o wartościach, a wspólne lutowanie może być okazją do budowania relacji, cierpliwości i zaufania. To propozycja dla młodych konstruktorów i ich opiekunów, którzy chcą, by technika była nie tylko użyteczna, lecz także mądra i bliska drugiemu człowiekowi.

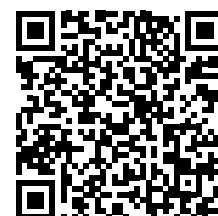
Luty to dobry moment, by wsłuchać się w sygnały, które do nas docierają – i te mierzalne, i te bardziej subtelne. Ten numer EdW prowadzi od analizy widma po migające serce LED, od fal radiowych po cyfrowe wzory życia. Bo w elektronice, podobnie jak w relacjach, najciekawsze rzeczy dzieją się wtedy, gdy technika spotyka się z uważnością.

Mariusz Ciszewski





pakiet promocyjny
KOCHAM SZACHY
7 e-booków z rabatem
50%



Dla prenumeratorów – 30% rabatu!

Promocja internetowa – w formularzu zamówienia online zaznacz pole „Jestem prenumeratorem wydawnictwa AVT, kupuję ze zniżką” i podaj swój numer prenumeraty.

W rubryce „Pocztą” zamieszczamy fragmenty listów od Czytelników. Szczególnie chętnie publikujemy komentarze do artykułów w bieżących wydaniach EdW oraz propozycje tematów artykułów, zadań i quizów.

Szyfry i kody w plikach Gerber

Szanowna Redakcja EdW,

mam 14 lat i od niedawna uczę się rysować płytki PCB w programie KiCad. Na razie dopiero zaczynam, ale udało mi się już zaprojektować swoje pierwsze płytki i nauczyć się, jak generować pliki Gerber, żeby można było zamówić płytkę w fabryce. Sam zamawiać ich jeszcze nie umiem, ale pomaga mi w tym znajomy mojego taty, który ma z tym doświadczenie.

Od niego dowiedziałem się, że każdy plik Gerber odpowiada za inną warstwę płytki. Jeden plik jest od miedzi, inny od tej białej farby z opisami, jeszcze inny od zielonej farby ochronnej, a kolejny od otworów. To było dla mnie bardzo ciekawe, bo wcześniej myślałem, że płytka to po prostu jeden rysunek.

Z ciekawości otwierałem te pliki Gerber w Notatniku i zobaczyłem tam jakieś dziwne znaki, liczby i kody, które wyglądają trochę jak szyfr. Zastanawiam się, jak to właściwie działa i dlaczego są tam same liczby i napisy, a nie na przykład zwykłe rysunki w formacie JPG albo BMP. Przecież każda z tych warstw to tak naprawdę jakiś rysunek.

Chciałbym zapytać, czy ktoś w Redakcji wie, o co w tym wszystkim chodzi i jak fabryka „czyta” takie pliki, żeby z nich zrobić prawdziwą płytkę. Jeśli tak, to byłoby super, gdybyście mogli to wytłumaczyć w prosty sposób. Szukałem informacji w polskim Internecie, ale niewiele udało mi się znaleźć, a temat wydaje mi się bardzo ciekawy.

Pozdrawiam serdecznie

Michał

Drogi Michale,

muszę przyznać, że Twój list sprawił mi prawdziwą radość – i to z bardzo osobistych powodów. Kiedy byłem w Twoim wieku, w elektronice najbardziej fascynowały mnie płytki drukowane. I wiesz co? Tak już zostało. Z płytkami drukowanymi związałem później swoją zawodową drogę i dziś, z uśmiechem, mówię o sobie, że jestem już chyba bardziej technologiem niż klasycznym elektronikiem.

Dlatego z ogromną sympatią czytałem Twoje rozważania o plikach Gerber – o tych „szyfrach”, liczbach i tajemniczych poleceniach, które w Notatniku zupełnie nie przypominają obrazków, a przecież rysują płytkę.

To bardzo trafna obserwacja i bardzo dobre pytanie.

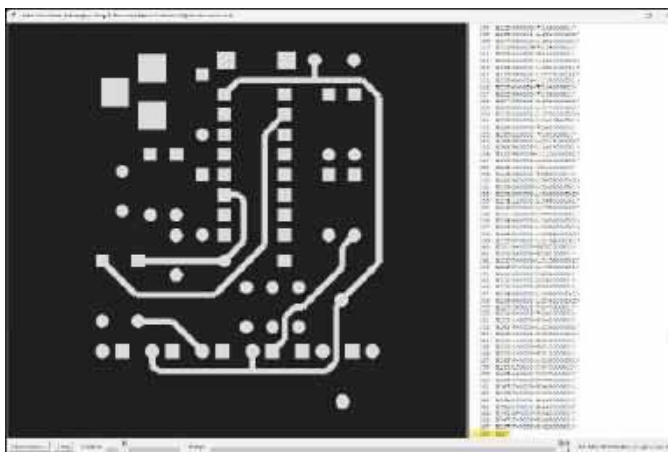
Na początek mała wskazówka: istnieje znakomita strona poświęcona wyłącznie temu tematowi – <https://www.artwork.com/gerber/index.htm>

To prawdziwa kopalnia wiedzy o formacie Gerber, jego strukturze i znaczeniu. Upprzedzam tylko, że materiały są w języku angielskim, ale zdecydowanie są warte uwagi.

Skąd wzięły się te „szyfry”?

Format Gerber ma swoje korzenie w czasach, gdy klisze do wykonywania płytek drukowanych powstawały przy użyciu fotoploterów – dużych, precyzyjnych maszyn optycznych. Na stronie <https://www.artwork.com/gerber/appl2.htm> znajdziesz nawet rysunek przedstawiający takie urządzenie.

Kluczowym elementem fotoplotera było koło apertur (ang. aperture wheel) – fizyczny



talerz z zestawem otworów o różnych kształtach i średnicach. Maszyna, rysując obraz warstwy, mechanicznie zmieniała aperturę, a następnie naświetlała światłoczułą kliszę: raz rysując linię, raz punkt, innym razem większy pad.

Plik Gerber był więc kodem sterującym maszyną: zmień aperturę, przesun się tu, naświetl, zgaś światło, przesun się dalej. Dzisiejsze technologie są zupełnie inne, ale struktura plików i te wszystkie „kody” pozostały – jako dziedzictwo tamtych czasów.

Małe marzenie z dzieciństwa...

I tu dochodzimy do czegoś, co szczególnie mnie poruszyło. Jako dzieciak marzyłem, żeby stworzyć program, który otwiera plik Gerber linia po linii i pokazuje, co dokładnie dzieje się z obrazem płytki. Żeby zobaczyć, jak z tych suchych poleceń krok po kroku wyłania się rysunek warstwy.

Nie byłem jednak na tyle biegłym programistą, by takie narzędzie napisać.

I nagle – po latach – do redakcji piśmenny Czytelnik, który swoim dociekliwym pytaniem przypomniał mi o tym marzeniu. Co więcej, zrobił to w czasach, gdy dostępna dla każdego jest sztuczna inteligencja, zdolna spełnić niemal każde techniczne życzenie. Wystarczy je dobrze opisać.

Tak też zrobiłem.

Po kilku podejściach powstała aplikacja w Pythonie, dokładnie taka, o jakiej marzyłem przed laty – a nawet lepsza. Aplikacja pozwala załadować plik Gerber i ma dwa panele:

- po lewej stronie rysuje się warstwa płytki,
- po prawej widać treść pliku Gerber, linia po linii.

Można nacisnąć przycisk Play i obserwować, jak po prawej żółta belka przesuwa się w dół po kolejnych liniach „szyfru”, a po lewej – dokładnie w tym samym rytmie – powstaje obraz płytki, linia po linii.

Prędkość rysowania da się regulować suwakiem, proces można zatrzymać, cofnąć się lub przejść do przodu. Można kliknąć dowolną linię i sprawdzić, co zostanie narysowane do tego momentu. Dodatkowo, najężdżając kursorem na linię, zobaczyć w dymku opis jej działania.

Z ogromną przyjemnością oddaję w Twoje (i pozostałych Czytelników) ręce narzędzie,

dzięki któremu możesz samodzielnie obserwować, co naprawdę dzieje się z rysunkiem warstwy Gerber na poziomie pojedynczych poleceń.

I na koniec – szczerze dziękuję. Dzięki Twojemu pytaniu mogłem, zupełnie niespodziewanie, spełnić jedną z własnych dziecięcych fantazji. To chyba jedna z najpiękniejszych rzeczy w naszej pracy.

Aplikację i przykładowy plik Gerber w formacie RS274X znajdziesz w materiałach do pobrania do bieżącego numeru EdW. Aby można było ją uruchomić (choćby dwuklikiem), wymagane jest wcześniejsze zainstalowanie środowiska Python. Najprościej i najbezpieczniej pobrać je z oficjalnej strony projektu: <https://www.python.org>

Serdecznie pozdrawiam

Redaktor naczelny

W artykule opisano podręczny generator i analizator sygnałów audio, oparty na tanim module Raspberry Pi Pico i paru innych elementach. Urządzenie oferuje funkcje generatora, oscyloskopu i analizatora widma. Może wytworzyć sygnał o przemiatanej częstotliwości i wykreślić charakterystykę częstotliwościową układu, potrafi też przeprowadzić analizę zawartości harmonicznych w celu sprawdzenia jakości sygnału. Analizator jest zasilany z akumulatora i mieści się w dłoni.



PICO – miniaturowy analizator elektroakustyczny

Analizator PICO to zwarte, poręczne urządzenie, zasilane z akumulatora. Wytwarza i analizuje podstawowe sygnały używane w elektroakustyce. Nadaje się do szeregu zadań, jak sprawdzanie wzmacniaczy, okablowania, filtrów i tym podobnych. Stanowi wygodne narzędzie do rozwiązywania problemów występujących w układach audio – w pracowni i w terenie. Można go podłączyć do płytki prototypowej i testować proste obwody, takie jak filtry RC.

Projekt został zainspirowany artykułem w Silicon Chip w dziale Circuit Notebook, opisującym analizator widma, w którym wykorzystano mikrokontroler dsPIC z wyświetlaczem LCD (sierpień 2023; www.siliconchip.au/Article/15908). Koncepcja analizatora PICO nawiązuje

również do naszego analizatora zniekształceń małej częstotliwości (kwiecień 2015; www.siliconchip.au/Article/8441).

Wykorzystaliśmy zawarte w tych projektach pomysły i dodaliśmy nowe funkcje. Jednym z interesujących zastosowań jest analiza zniekształceń napięcia sieci energetycznej 230 V, którego kształt powinien być sinusoidalny – ale niekiedy mało przypomina sinusoidę! Aby dokonać takiej analizy, należy dołączyć uzwojenie wtórne niemal dowolnego transformatora sieciowego do wejścia analizatora i przełączyć go w tryb analizy zniekształceń.

Podobnie jak konstrukcje wspomniane wcześniej, analizator PICO stosuje do badania częstotliwości składowych sygnału transformatę Fouriera. Wyniki

transformaty można wyświetlić w postaci wykresu widma. Można też przeprowadzić analizę z przemiataniem częstotliwości. W artykule w Silicon Chip z kwietnia 2015 roku szczegółowo wyjaśniono zastosowanie transformaty Fouriera oraz sposób wykorzystania jej wyników do pomiaru zniekształceń.

Projekt

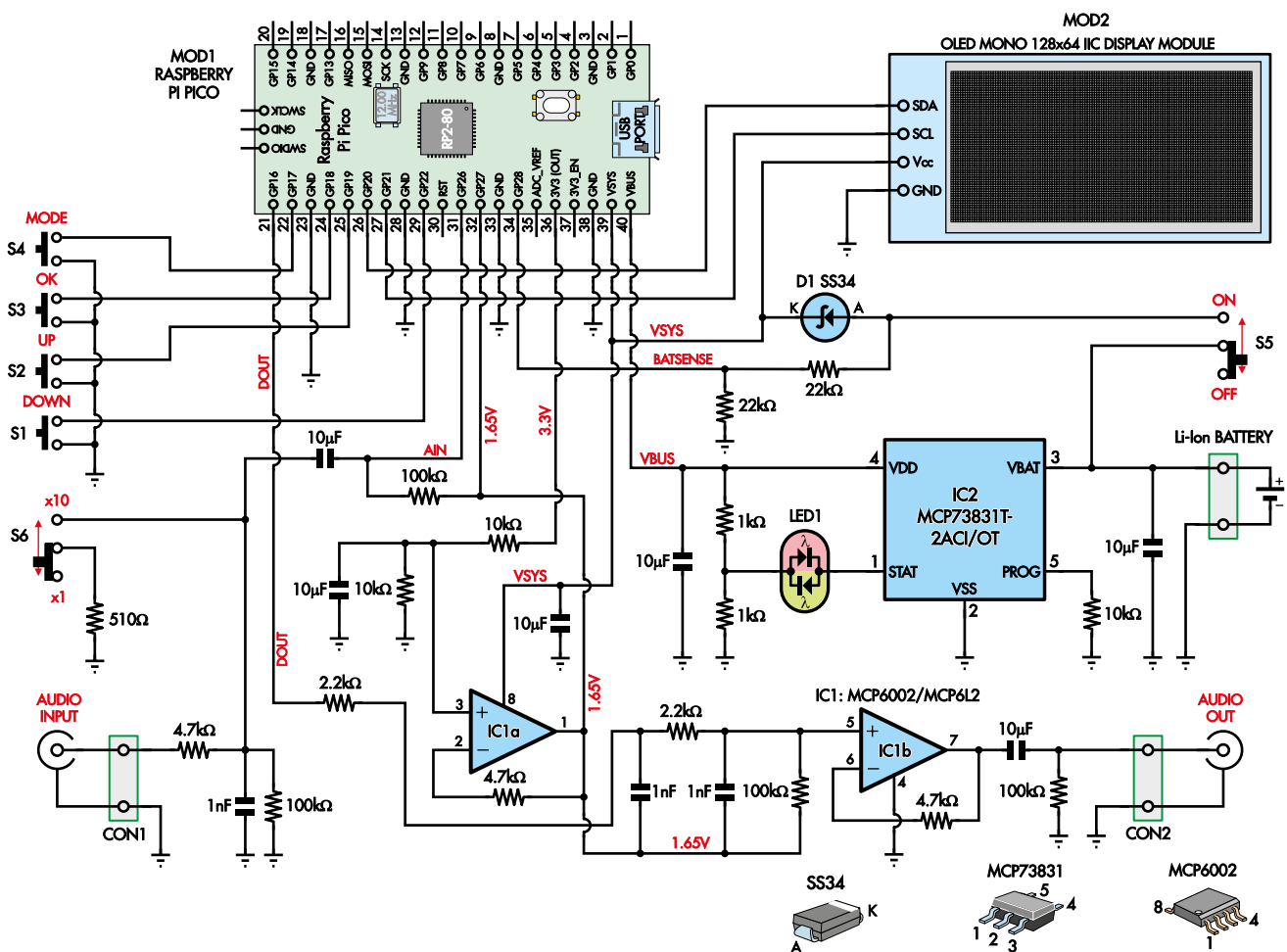
Projektując opisywany układ mieliśmy na uwadze, że powinien być on miniaturowy i tani. Cała konstrukcja mieści się w najmniejszej obudowie Altronics z serii Jiffy (UB5) o wymiarach zaledwie 83×54 mm. Panel przedni to spodnia strona głównej płytki drukowanej, wpuszczonej w górną część obudowy. Wysokość urządzenia wynosi zaledwie 28 mm, czyli sporo mniej niż w przypadku użycia dołączonej pokrywki.

Zastosowano wyświetlacz typu OLED o przekątnej 1,3 cala (33 mm) – najmniejszy typ wyświetlacza zdolnego do wyświetlania grafiki (przypis redaktora: rozpowszechnione są również mniejsze typy OLED, o przekątnej 0,91 cala). Można wyświetlać kilka wierszy tekstu. Tego rodzaju wyświetlacza użyliśmy w pęsetce pomiarowej „SMD Tweezers – kolejne wcielenie pęsety pomiarowej SMD” (EdW 10–11/2025).

W projekcie nie zastosowano drogich przetworników analogowo-cyfrowych (ADC) ani cyfrowo-analogowych (DAC). Do próbkowania sygnału wejściowego wykorzystano wewnętrzny 12-bitowy przetwornik ADC mikrokontrolera (patrz Wady wewnętrznego przetwornika ADC),

Specyfikacja

- Generator sygnału akustycznego o amplitudzie do 3 Vpp (1,06 Vsk) o nastawianej częstotliwości
- Kształt przebiegu: sinusoidalny, prostokątny, trójkątny, piłokształtny i biały szum
- Wejście sygnału audio z przełączanymi zakresami 3,6 Vpp/34 Vpp (1,27 Vsk/12 Vsk)
- Tryby oscyloskopu i analizatora widma
- Analiza harmoniczna z pomiarem współczynnika THD od 0,3% (1,2 Vsk, 1,2 kHz)
- Pomiar i monitorowanie zniekształceń napięcia sieci energetycznej (z zewnętrznym transformatorem)
- „Przemiatanie” częstotliwości i kreślenie charakterystyki częstotliwościowej
- Gniazda wejściowe i wyjściowe RCA
- Zasilanie z USB lub z wewnętrznego akumulatora
- Wyświetlacz OLED o rozdzielczości 128×64 px i przyciski sterujące
- Mały i przenośny
- Możliwość sterowania przez wirtualny port szeregowy USB
- Pobór prądu typowo około 50 mA
- Z naładowanym akumulatorem 600 mAh działa przez około 12 godzin



Rysunek 1. Analizator to przede wszystkim oprogramowanie pracujące na module Pico. Moduł można było zmontować na płytce krawędzią, co oszczędziło sporo miejsca na płytce drukowanej

a do przetwarzania DAC użyto metody PWM (modulacja szerokości impulsu).

Moduł Raspberry Pi Pico zawiera układ scalony impulsowego stabilizatora napięcia 3,3 V. Układ może działać w trybie PFM (modulacja częstotliwości impulsów) lub PWM. Używamy trybu PWM, ponieważ przy małym obciążeniu prądowym tryb PFM może wprowadzać tętnienia o niskich częstotliwościach, trudne do odfiltrowania.

W prezentowanym układzie istotne jest czyste napięcie zasilania 3,3 V, ponieważ napięcie to jest używane jako poziom odniesienia dla ADC, wyznacza również poziom wyjściowy wytwarzanych przebiegów. Stabilizator impulsowy zasadniczo nie jest w takim przypadku zalecany. Użycie stabilizatora wbudowanego w moduł Raspberry Pi wyeliminowało jednak konieczność zastosowania dodatkowego układu.

Szczegóły układowe

Rysunek 1 przedstawia schemat analizatora. MOD1 to moduł mikrokontrolera Raspberry Pi Pico. Stopień

wejściowy analizatora otrzymuje sygnał przez gniazdo RCA dołączone do złącza CON1. Rezystor 4,7 kΩ w połączeniu z kondensatorem 1 nF tworzy filtr dolnoprzepustowy o częstotliwości granicznej (−3 dB) około 34 kHz. Filtr tłumi składowe o wielkich częstotliwościach, które w przetworniku ADC mogłyby powodować aliasing ([nakładanie się widm; przypis redaktora](#)).

Gdy do gniazda nie jest nic podłączone, rezystor 100 kΩ utrzymuje w węzle wejściowym potencjał masy. Kondensator 10 μF sprzęga zmiennoprądowo sygnał z wejściem analogowym modułu Pico. Wejście (AIN, pin 31) jest polaryzowane stałym potencjałem 1,65 V (połowa napięcia zasilania 3,3 V) przez inny rezystor 100 kΩ.

Obwód wejściowy nieznacznie tłumi dźwięk – do około 91% jego pierwotnego poziomu. Oznacza to, że można bez przesterowania mierzyć napięcie do 3,6 Vpp (międzyzszczytowe), co odpowiada około 1,2 Vsk (wartość skuteczna).

Przełącznik S6 może równolegle do pierwszego rezystora 100 kΩ dołączyć rezystor

510 Ω, zmieniając parametry dzielnika z rezystorem 4,7 kΩ. Pozwala to na pomiar sygnałów o poziomach do 34 Vpp (12 Vsk) bez obcinania. Rezystor ten zmienia jednak charakterystykę filtra i przepuszcza więcej składowych o dużej częstotliwości niż na zakresie ×1.

Sygnał wyjściowy z modułu Pico (DOUT, pin 21) to przebieg PWM o częstotliwości około 250 kHz. Przebieg ten najpierw przechodzi przez dwustopniowy filtr RC. Każdy stopień składa się z rezystora 2,2 kΩ i kondensatora 1 nF. Kolejny rezystor 100 kΩ polaryzuje wyjście filtra potencjałem 1,65 V, który wyznacza stałą sygnału. Dwa stopnie RC filtru dają podobną częstotliwość graniczną jak filtr w stopniu wejściowym, tłumiąc artefakty sygnału PWM o częstotliwości 250 kHz łącznie o 24 dB (12 dB na każdy stopień). Rezultat nie spełnia może norm Hi-Fi, ale dla naszych celów jest wystarczający. Sygnał jest buforowany przez wtórnik IC1b, a następnie dołączony zmiennoprądowo do złącza CON2 przez kondensator 10 μF. Rezystor 100 kΩ polaryzuje wyjście potencjałem masy.

Filtrowanie i polaryzacja sprawiają, że napięcie wyjściowe wynosi około 3,1 Vpp (około 1,1 Vsk). Większa wartość nie byłaby możliwa z racji ograniczeń wnoszonych przez wzmacniacz operacyjny.

Do cyfrowych pinów wejściowych modułu Pico są dołączone przyciski S1...S4. Służą one do sterowania interfejsem użytkownika. Gdy przyciski są otwarte, wewnętrzne rezystory podciągające w module utrzymują odpowiednie piny w stanie wysokim. Naciśnięcie przycisku ściąga pin do masy, wymuszając stan niski.

Zasilanie

Wewnątrz modułu MOD1 znajduje się dioda Schottky'ego, włączona od pinu VBUS (40) do pinu VSYS (39). Z napięcia VSYS jest na module Pico zasilany stabilizator impulsowy, z wyjściem na pinie 3V3. Zauważmy, że moduł Pico jest dołączony do układu analizatora tylko jedną stroną. Można więc było go zamontować na płycie drukowanej pionowo, oszczędzając dużo miejsca.

Cały układ może być zasilany ze złącza USB na module Pico. Wówczas na pinie VSYS panuje napięcie około 4,7 V. Inna możliwość to zasilanie z akumulatora litowego, podawane przez diodę D1 gdy przełącznik S5 jest zamknięty, co daje w węźle VSYS około 3,4...3,9 V. Stabilizator na module Pico działa przy napięciu wejściowym 1,8 V... 5,5 V, więc wymienione napięcia VSYS mieszczą się w jego zakresie roboczym.

Gdy zasilanie przychodzi z USB, przez układ IC2 jest ładowany akumulator. Dwa kondensatory 10 µF zapewniają wymagane odprężanie wejścia i wyjścia układu. Rezystor 10 kΩ między pinem PROG (5) a masą ustala prąd ładowania na 100 mA. Pin STAT (1) jest podczas ładowania w stanie niskim, a po zakończeniu ładowania przechodzi w stan wysoki. Stan ładowania wskazuje dwukolorowa dioda LED: czerwony – ładowanie w toku, zielony – ładowanie zakończone. Prąd każdej sekcji LED ograniczają oddzielne rezystory 1 kΩ.

MOD2 to moduł wyświetlacza OLED o przekątnej 1,3 cala (33 mm). Jego pin VCC jest zasilany napięciem na pinie VSYS modułu Pico. Wbudowany w wyświetlacz stabilizator daje napięcie 3,3 V do zasilania i dla wewnętrznych rezystorów podciągających na liniach SDA i SCL interfejsu I²C. Linie te są sterowane z odpowiednich pinów modułu Pico.

IC1 jest podwójnym niskonapięciowym wzmacniaczem operacyjnym, również zasilanym z węzła VSYS, z dodanym kondensatorem filtrującym 10 µF. Napięcie VSYS jest zawsze wyższe niż 3,3 V, więc jest zapewniony

większy zapas napięcia niż gdybyśmy zasilali wzmacniacz z linii 3V3.

Napięcie linii 3V3, dzielone na pół przez parę rezystorów 10 kΩ i filtrowane przez kondensator 10 µF, służy jako napięcie odniesienia 1,65 V. Jest ono buforowane przez wzmacniacz operacyjny IC1a o jednostkowym wzmocnieniu.

Moduł Pico ma cztery kanały ADC, z których jeden jest wewnętrznie podłączony do VSYS za pośrednictwem dzielnika, a drugiego używamy do wprowadzania sygnału audio. Pozostają więc dwa. Dołączyliśmy jeden z nich do napięcia odniesienia 1,65 V, aby moduł Pico mógł sprawdzać, czy napięcie to jest prawidłowe. **Przypis redaktora: chodzi tu nie tyle o pomiar bezwzględnej wartości tego napięcia, ile o potwierdzenie, że jest ono połową napięcia odniesienia przetwornika ADC.** Cztery kanały ADC mierzy napięcie dzielnika na dwóch rezystorach 22 kΩ, dołączonego do akumulatora za przełącznikiem S5. Pozwala to na pomiar napięcia akumulatora – ale tylko gdy S5 jest zamknięty, co gwarantuje, że akumulator nie będzie rozładowywany po wyłączeniu urządzenia.

Oprogramowanie

Do utworzenia oprogramowania użyliśmy środowiska Arduino IDE, głównie dlatego, że dostępne są w nim liczne biblioteki. Używamy bibliotek OLED od Adafruit, które ułatwiają tworzenie grafiki do sporządzania diagramów widma, oscylogramów i wykresów odpowiedzi częstotliwościowej.

Oprogramowanie do wytwarzania dźwięku to dość prosta implementacja metody PWM, w której dla każdej nowej próbki jest aktualizowane wypełnienie przebiegu, co zapewnia zmienny kształt fali. Implementacja jest oparta na oprogramowaniu, które napisaliśmy dla Pico Backpack z akustycznym wyjściem stereo (www.siliconchip.au/Article/15236). Używamy przetwarzania PWM ośmio-bitowego, chociaż próbki są obliczane i przechowywane w formacie 16-bitowym.



Na prawym boku obudowy są dwa otwory na gniazda RCA oraz wycięcie na przełącznik S6

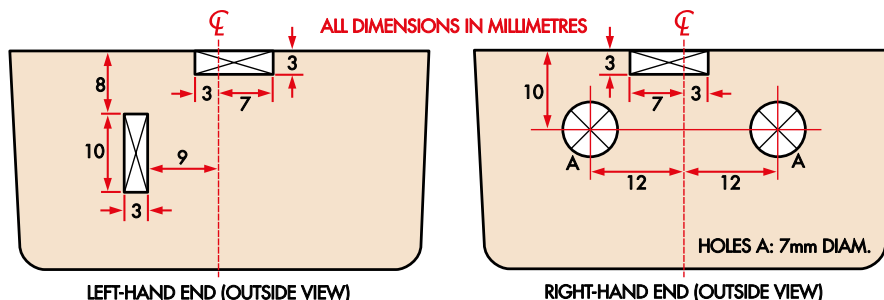
Wartości PWM pochodzą z górnych ośmiu bitów. „Reszta” z dolnych ośmiu bitów służy do zrealizowania ditheringu (rozpraszania) wartości przez kilka okresów PWM w każdym okresie próbkowania, co nieznacznie poprawia wynikową rozdzielczość.

Dane dla przetwornika PWM pochodzą z bloku próbek o długości około 200 ms. Oznacza to, że generowana częstotliwość nie musi być całkowitą podwielokrotnością częstotliwości próbkowania, ponieważ blok zawiera kilka okresów sygnału. Nie dotyczy to jedynie najniższych częstotliwości.

Do obliczania i aktualizowania próbek oraz dodawania ditheringu używamy drugiego rdzenia procesora. Rdzeń ten nie ma w zasadzie innych zadań, co zapewnia niezakłócone wytwarzanie fali dźwiękowej.

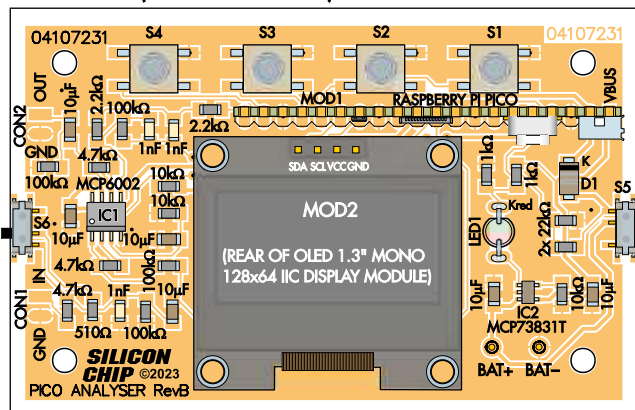
Podobny sposób zastosowano do próbkowania wejścia analogowego. 12-bitowy przetwornik ADC działa z częstotliwością 490 kHz, zbliżoną do maksymalnej 500 kHz. Blok DMA zapisuje blok próbek przez około 1/10 sekundy bez przerywania pracy procesora. Sprawia to, że można wykrywać częstotliwości począwszy od około 10 Hz.

Jakość przetwornika ADC nieco nas rozczarowała. Jak się okazało, w układzie RP2040 na module Pico są pewne problemy z blokiem ADC (patrz *Wady wewnętrznej przetwornika ADC*). Aby nieco zrekomensować te braki, nasze oprogramowanie

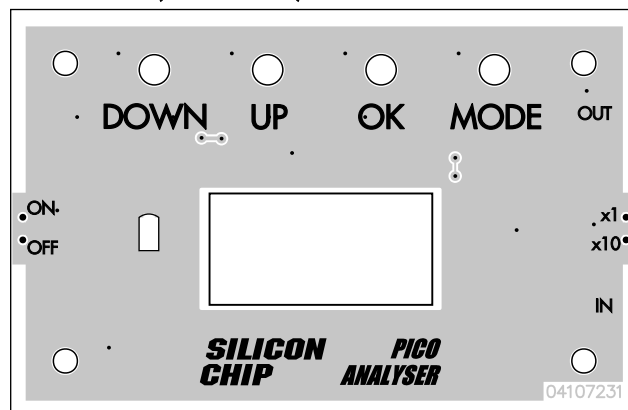


Rysunek 2. Część tych otworów można wykonać bez mierzenia. Miejsca wycięć na przełączniki w górnej części obudowy można zaznaczyć używając jako szablonu płytki drukowanej. Otwory na gniazda RCA nie muszą być umieszczone precyzyjnie, ponieważ gniazda są do płytki dołączone przewodami

REAR OF PCB (COMPONENT SIDE)



FRONT OF PCB (FRONT PANEL)



Rysunek 3. Spód płytki drukowanej jest jednocześnie panelem przednim urządzenia. Wszystkie elementy należy zamontować powierzchniowo. Jest trochę ciasno, bo wyświetlacz i moduł Pico są blisko siebie. Zalecamy najpierw zamontować moduł Pico i sprawdzić, czy jest wyśrodkowany z otworem w obudowie na gniazdo USB, a następnie zamontować wyświetlacz i od razu sprawdzić jego działanie. Pamiętajmy, że wyświetlacz „patrzy” w stronę spodu płytki

do odczytów wyników ADC stosuje korekty. Trochę to pomaga, ale efektywna rozdzielczość przetwornika wynosi zaledwie około dziewięciu bitów.

W trybie oscyloskopu wyświetlane są wartości sygnału nieobrobionego. Oś pionowa wykresu ma na wyświetlaczu rozdzielczość 50 pikseli. W pozostałych trybach sygnał jest decymowany ([próbkowany z częstotliwością próbkowania mniejszą niż oryginalna; przypis redaktora](#)) i przeprowadzana jest transformata Fouriera mająca na celu znalezienie częstotliwości składowych próbkowanego przebiegu. Spora część oprogramowania jest zaangażowana w realizację różnych procedur graficznych i stron interfejsu użytkownika.

Konstrukcja

Najpierw, używając nieobłożonej płytki drukowanej, zaznaczamy jej kontur na obudowie. Następnie wykonujemy wycięcia w lewym i prawym boku obudowy, pokazane na **rysunku 2**. Jeden bok ma wycięcie na przełącznik S5 i gniazdo USB Pico, a drugi – na przełącznik S6 oraz dwa otwory na gniazda RCA. Miejsca wycięć można zaznaczyć przy pomocy krawędzi płytki drukowanej, co będzie łatwiejsze niż znalezienie punktu środkowego linijką.

Oстрым nożem lub pilką do metalu wykonujemy parę pionowych nacięć z każdej strony na niepełną głębokość. Kolejne nacięcia wykonujemy ostrym nożem wzdłuż dolnej krawędzi. Następnie ostrożnie wyginamy usuwany fragment obudowy. Powinien się odłamać wzdłuż naciętej linii. W razie potrzeby wyrównujemy rogi i krawędzie do odpowiedniej głębokości małym pilnikiem lub ostrym nożem.

Zaznaczamy miejsce na gniazdo USB. Zaczynamy od wywiercenia dwóch lub trzech

otworów wewnątrz linii. Następnie wyrównujemy krawędzie gniazda małym pilnikiem lub ostro zakończonym nożem. Możemy obrobić tę szczelinę precyzyjnie i później użyć jej, aby dokładnie umieścić w obudowie płytkę analizatora. Jeśli nie damy rady zapewnić odpowiedniej precyzji, lepiej będzie szczelinę trochę poszerzyć.

Gniazda RCA montujemy w wywierconych otworach. Można je wykonać wiertłem krętym lub stopniowym. Ich dokładne położenie nie jest krytyczne, ponieważ gniazda są łączone z płytką przewodami. Przedstawione wymiary dobrze się sprawdziły w prototypie. Wiercenie proponujemy najpierw przeprowadzić wiertłem pilotażowym 3 mm, co ułatwi wyrównanie położenia otworów. Dla użytych przez nas gniazd RCA średnice otworów końcowych wynoszą 7 mm. Sprawdźcie, czy dla gniazd używanych przez Was nie będzie odpowiednia inna średnica.

Montaż płytki

Większość elementów to standardowe, dość duże podzespoły SMD. Kilka elementów musi być zamontowanych w nieco niekonwencjonalny sposób. Zalecamy rozpoczęcie montażu od elementów SMD. Będzie potrzebna lutownica z cienkim grotem, cienka cyna, pasta topnikowa, pęseta i dobre oświetlenie. Przyda się również plecionka lutownicza oraz rozpuszczalnik do usuwania nadmiaru topnika.

Aby się nie narażać na działanie oparów z topnika, należy stosować system wentylacyjny (np. okap, lub odciąg oparów lutowniczych). Jeśli go nie ma, należy pracować na świeżym powietrzu.

Rysunek 3 przedstawia rozmieszczenie i orientację elementów, a także zdjęcie przedstawiające płytkę drukowaną z elementami SMD.

Elementy są rozmieszczone dość blisko siebie. Najmniejszym podzespołem jest IC2, więc montaż najlepiej zacząć od niego. Nakładamy trochę topnika w postaci pasty na pady, kładziemy na nich i wyrównujemy IC2 (będzie pasować tylko na jeden sposób). Po stronie z dwoma wyprowadzeniami lutujemy jedno z nich, po czym sprawdzamy, czy wszystkie pozostałe piny znajdują się na właściwych miejscach. W razie potrzeby wyrównujemy położenie układu. Lutujemy pozostałe piny, a następnie poprawiamy lutowanie pierwszego piny.

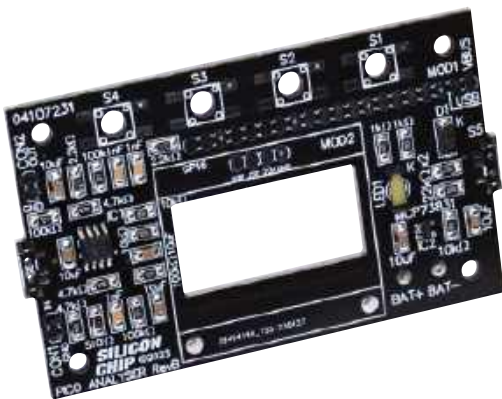
Sprawdzamy, czy nie ma zwarc. Jeśli są, lub gdy jest nadmiar cyny, usuwamy ją plecionką z dodatkiem topnika.

W podobny sposób przylutowujemy IC1. Jego piny mają szerszy rozstaw, więc lutowanie powinno być łatwiejsze. Upewniamy się, że pin 1 tego układu (może być oznaczony paskiem wzdłuż jednej z krawędzi) jest przy kropce na warstwie opisowej płytki drukowanej.

Na całej płytce występują kondensatory o dwóch wartościach. Należy uważać, aby ich nie pomylić ze sobą. Typy 10 µF będą prawdopodobnie grubsze niż 1 nF. Również rezystory muszą być umieszczone we właściwych miejscach. Do elementów biernych należy użyć tej samej podstawowej techniki lutowania:



Spodnia strona płytki drukowanej czyli panel czołowy



To zdjęcie może posłużyć jako wzór do montażu małych elementów. Na tym etapie montażu dobrze jest usunąć nadmiar topnika i dopiero potem przejść do umieszczenia elementów końcowych: przełączników, diody LED, modułu Pico i wyświetlacza

lutujemy jeden koniec elementu, sprawdzamy i ewentualnie poprawiamy położenie elementu, a następnie lutujemy drugi koniec.

Dioda stanowi większy element SMD. Przy montażu istotna jest prawidłowa biegunowość – paskiem katody w kierunku litery „K” na płytce drukowanej. Gdyby biegunowość została odwrócona, to mogłoby się zdarzyć, że bateria zostałaby dołączona bezpośrednio do zasilania USB, co prawdopodobnie spowodowałoby uszkodzenie układu.

Ten etap montażu to dobry moment na wy-czyszczanie płytki rozpuszczalnikami topnika. Przeprowadzenie czyszczenia na tym etapie zapobiegnie przedostaniu się rozpuszczalnika do wnętrza przełączników. W roli rozpuszczalnika dobrze się sprawdzi alkohol izopropylowy. Zanim przejdziemy do dalszej części montażu, musimy odczekać, aż płytka dokładnie wyschnie.

Montujemy przełączniki suwakowe S5 i S6. Mają one krótkie końcówki, ale są łatwe w montażu, ponieważ w ich dolnej części znajdują się kołki pozycjonujące, które wpasowują się w otwory w płytce drukowanej.

Lutujemy jedno wyprowadzenie, sprawdzamy, czy przełącznik leży płasko, a następnie lutujemy pozostałe końcówki.

Dalej montujemy cztery odwrócone przyciski S1...S4 (przypis redaktora: chodzi o typy, w których trzpień przycisku jest po innej stronie płytki niż wyprowadzenia). Dobrze będzie wygiąć ich wyprowadzenia, tak aby trzpień przełączników głębiej wystawał przez otwory w płytce drukowanej. Również w tym wypadku warto przylutować najpierw jedno wyprowadzenie i sprawdzić, czy trzpień jest dokładnie w środku otworu. Wyeliminuje to możliwość zakleszczenia się trzpienia w płytce drukowanej, i zapewni estetyczny wygląd panelu przedniego. Gdy przycisk jest już umieszczony prawidłowo, lutujemy pozostałe wyprowadzenia. Używamy sporej ilości lutowia, aby zapewnić dobrą wytrzymałość mechaniczną.

Elementy specjalne

Dioda LED jest zamontowana nietypowo. Dwukolorowe LED-y SMD mają zwykle niezależne wyprowadzenia każdej sekcji, co sprawia, że pady są małe i trudno do nich cokolwiek przylutować. Użyliśmy więc LED-a prze-wlekane go o średnicy 3 mm, zamontowanego powierzchniowo i skierowanego odwrotnie (świecącego poprzez płytkę). Oznaczenie „Kred” odpowiada katodzie sekcji czerwonej. Jeśli nie jesteśmy w stanie ustalić, które to wyprowadzenie, zamontujcie LED w dowolny sposób i potem ewentualnie zmieńcie jego orientację. Ostrożnie wyginamy wyprowadzenia o 180° i przycinamy je tak, aby były nieco dłuższe niż soczewka diody. Jak widać na zdjęciu, soczewka LED-a jest skierowana w stronę otworu w masce lutowniczej (w kierunku płytki drukowanej). Wyprowadzenia LED-a lutujemy do padów, posługując się pęsetą. Lutujemy najpierw jedno wyprowadzenie, poprawiamy pozycję LED-a, lutujemy drugą końcówkę i na końcu poprawiamy lutowanie pierwszej.

Następnym elementem jest moduł Pico. Przed przystąpieniem do montażu należy sprawdzić, czy płytka drukowana (z zamontowanymi przełącznikami S5 i S6) znajduje się w jednej płaszczyźnie z wycięciami w obudowie. Górna część płytki drukowanej powinna znajdować się na poziomie otaczających ją boków obudowy. Ma to na celu zapewnienie prawidłowej pozycji złącza USB na module Pico względem otworu z boku obudowy. Może się okazać, że trzeba będzie skorygować wycięcia.

Umieszczamy moduł Pico na płytce drukowanej pod kątem prostym. Zwracamy uwagę na jego orientację; pin VBUS powinien być blisko krawędzi płytki drukowanej, a GP16 – z drugiej strony. Gniazdo USB powinno znaleźć się ponad odpowiednim oznaczeniem na warstwie opisowej płytki. Przylutowujemy najpierw tylko jedno ze skrajnych wyprowadzeń modułu. Moduł powinien być w niewielkiej odległości od brzegu płytki, a złącze USB powinno lekko wystawać. Ostrożnie pozycjonujemy moduł Pico, aby był pod kątem prostym do płytki drukowanej. Wkładamy płytkę do obudowy i sprawdzamy, czy gniazdo USB jest równo z otworem w boku. Pamiętajmy, że górna część płytki drukowanej będzie znajdować się w jednej płaszczyźnie z górną częścią obudowy.

Gdy położenie modułu Pico jest prawidłowe, lutujemy jego pozostałe wyprowadzenia. Najłatwiej będzie lutować moduł po stronie przełączników, przykładając lutownicę do padów na module. Lutowie na każdym pinie powinno tworzyć obfite zaokrąglenie, aby moduł był solidnie umocowany.

Teraz przycinamy cztery kawałki cienkiego drutu o długości po około 1 cm. Najlepiej aby każdy był odrobinę innej długości, co ułatwi wsuwanie do nich padów wyświetlacza MOD2. Jako drutu można użyć odciętych wyprowadzeń LED. Przytrzymując każdy odcinek drutu pęsetą, lutujemy je w pozycji pionowej do padów MOD2 na płytce. Z wyświetlacza OLED zdejmujemy folię ochronną i umieszczamy wyświetlacz ekranem do dołu ponad drutami, równoległe do płytki drukowanej. „Lustrzane” oznaczenia wyprowadzeń na płytce drukowanej będą odpowiadać oznaczeniom na wyświetlaczu. Teraz przylutowujemy każdy drut do wyświetlacza. Zachowujemy ostrożność, ponieważ połączenie każdego drutu z płytką drukowaną nagrzej się, a drut będzie utrzymywany tylko napięciem powierzchniowym lutowia. Może być konieczne przytrzymanie drutu pęsetą. Sprawdzamy, czy wyświetlacz jest dokładnie wyśrodkowany względem padów do oznaczeń na warstwie opisowej. Jeśli nie, będzie to widoczne podczas użytkowania.

Nadszedł teraz czas na przeprowadzenie paru szybkich testów, aby upewnić się, że moduł



W pełni okablowany analizator z otwartą pokrywą. Zwróćmy uwagę na sposób montażu diody LED, modułu Pico i wyświetlacza. Otwarcie pokrywy jest możliwe po wyciągnięciu przewodów z pojemnika akumulatora, co widać na rysunku 1. Warto zabezpieczyć i zaizolować przewody akumulatora dużą ilością silikonu o neutralnym utwardzaniu

Pico i wyświetlacz są prawidłowo przylutowane. W tym celu moduł Pico będzie musiał zostać zaprogramowany.

Programowanie modułu Pico

Trzymając wciśnięty biały przycisk BOOTSEL na module Pico, podłączamy moduł kablem USB do komputera. Jeśli moduł nie był nigdy programowany, przytrzymywanie przycisku może nie być konieczne. Moduły Pico sprzedawane w formie zestawów zazwyczaj są niezaprogramowane, ponieważ ułatwia to życie konstruktorom.

Kopiujemy plik 0410723A.UF2 do wirtualnego napędu dyskowego RPI-RP2, który powinien pojawić się w systemie plików komputera. Jeśli wszystko jest w porządku, wyświetlacz OLED powinien się zaświecić po około sekundzie. Jeśli tak nie jest, sprawdzamy połączenia lutowane i rozmieszczenie elementów na płytce.

Sprawdzamy czy wyświetlacz leży symetrycznie w wycięciu w płytce drukowanej. Jeśli nie, możemy go ostrożnie przemieścić.

Montaż końcowy

Po wypozycjonowaniu wyświetlacza mocujemy go dwoma odcinkami drutu przez dolne otwory do odpowiednich padów na płytce drukowanej. Postępujemy podobnie jak w przypadku czterech połączeń na górze wyświetlacza.

Przygotowujemy gniazda RCA. Przycinamy dwa odcinki białego przewodu i dwa odcinki czarnego, wszystkie o długości około 4 cm. Kolory nie są istotne, ale powinny być kontrastujące, co ułatwi ich identyfikację. Przylutowujemy jeden koniec każdego białego przewodu do środkowego zacisku gniazda RCA, a końce czarnych przewodów do podkładek, które będą połączeniem masy.

Montujemy gniazda w otworach w obudowie, zabezpieczając podkładki nakrętkami. Gniazda przekręcamy tak, aby przewody wystawały na zewnątrz z górnej części obudowy. Przewody zaginamy następnie wokół brzegu obudowy.

Umieszczamy płytkę drukowaną do góry nogami obok obudowy i przylutowujemy przewody, jak na zdjęciach. Dwa przewody czarne idą do padów GND na CON1 i CON2, a przewody białe – do odpowiednich padów oznaczonych IN i OUT. Aby zapewnić solidne połączenie, dajemy sporo lutowia.

Przymocowujemy pojemnik akumulatora do dolnego tylnego rogu obudowy, z wlotem skierowanym na zewnątrz. Stosujemy silikon o neutralnym utwardzaniu lub podobny klej wypełniający szczeliny. Przewody pojemnika przylutowujemy do pól BAT+ i BAT- na płytce drukowanej (przewód czerwony do BAT+). W prototypie przedłużyliśmy

Wykaz elementów:

- 1 dwustronna płytka drukowana; nr Silicon Chip 04107231; 83 × 50 mm, z czarną soldermaską
- 1 obudowa UB5 Jiffy (83 × 53 × 30 mm)
- 2 gniazda RCA do montażu w obudowie (CON1, CON2) [Altronics P0161]
- 1 pojemnik baterii na pojedyncze ogniwo AA, z przewodami
- 1 akumulator litowo-jonowy 14500, rozmiar AA
- 1 moduł Raspberry Pi Pico, zaprogramowany wsadem 0410723 A.UF2 (MOD1)
- 1 wyświetlacz OLED 1,3 cala (33 mm) (MOD2) [Silicon Chip SC5026]
- 4 przyciski SMD do montażu odwrotnego (S1...S4) [Adafruit 5410]
- 2 przełączniki suwakowe SPDT SMD (S5, S6)
- 4 podkładki M3 o grubości 1,5 mm
- 2 odcinki przewodu izolowanego o długości po 20 cm (np. biały i czarny)
- 1 odcinek drutu nieizolowanego o długości 4 cm (np. z odciętych wyprowadzeń diody LED1)
- 1 mała tubka uszczelniająca silikonowego o neutralnym utwardzaniu
- 1 krótki kabel RCA-RCA (do testowania i kalibracji)

Półprzewodniki:

- 1 podwójny wzmacniacz operacyjny „rail-to-rail” MCP6002 lub MCP6L2, obudowa SOIC-8 (IC1)
- 1 regulator ładowania akumulatora Li-ion MCP73831-2 ACI/OT, obudowa SOT-23-5 (IC2)
- 1 dwukolorowa (czerwona/zielona) dioda LED, średnica 3 mm (LED1)
- 1 dioda Schottky'ego 40 V/3 A SS34, obudowa DO-214 (D1)

Kondensatory: (wszystkie w rozmiarze M3216/1206, ceramiczne, X7R)

- 6 szt. 10 µF/16 V lub więcej
- 3 szt. 1 nF/50 V

Rezystory: (wszystkie w rozmiarze M3216/1206, 1%, 1/8 W)

- 4 szt. 100 kΩ 2 szt. 2,2 kΩ 2 szt. 22 kΩ 2 szt. 1 kΩ
- 3 szt. 10 kΩ 1 szt. 510 Ω 3 szt. 4,7 kΩ

Zestaw do samodzielnego montażu analizatora PICO

Zestaw SC6772 zawiera płytkę drukowaną i wszystkie elementy bezpośrednio na niej montowane. Moduł Pico jest dostarczany niezaprogramowany. Użytkownik musi go zaprogramować z komputera przez kabel USB.

niedługo jeden z przewodów baterii, co umożliwiło całkowite wyjmowanie płytki drukowanej z obudowy, ułatwiając testowanie i montaż. Przewody przylutowujemy dużą ilością lutowia, podobnie jak w przypadku gniazd RCA. Aby dodatkowo zabezpieczyć przewody i zaizolować gołe końce, wokół padów BAT+ i BAT- nakładamy silikon. Nie można dopuścić, aby przewody te się poluzowały, ponieważ spowodują wtedy z dużym prawdopodobieństwem zwarcie akumulatora.

Trochę silikonu można również nanieść na złącza CON1 i CON2 na płytce drukowanej, co zabezpieczy połączenia audio, a także na odsłonięte elementy metalowe pojemnika akumulatora, na przykład dookoła styków baterii.

Kiedy silikon ulegnie całkowitemu utwardzeniu, do pojemnika wstawiamy akumulator, uważając na biegunowość. Włączamy analizator wyłącznikiem S5 i sprawdzamy, czy po około sekundzie zaświeci się wyświetlacz. Jeśli nie, akumulator wyjmujemy i sprawdzamy, co może być przyczyną problemów.

Do gniazda USB podłączamy zasilanie i sprawdzamy, czy dioda LED świeci się początkowo na czerwono, a następnie, gdy akumulator skończy się ładować, na zielono. Jeśli LED początkowo świeci na zielono, może to oznaczać, że jest on wmontowany odwrotnie.

Przed dokonywaniem jakichkolwiek zmian w układzie zawsze dobrze jest wyjąć akumulator.

Płytkę drukowaną jest mocowana śrubkami do czterech słupków obudowy. Należy uważać,

aby podczas wkręcania śrubek nie przygnieść żadnych przewodów.

Kalibracja

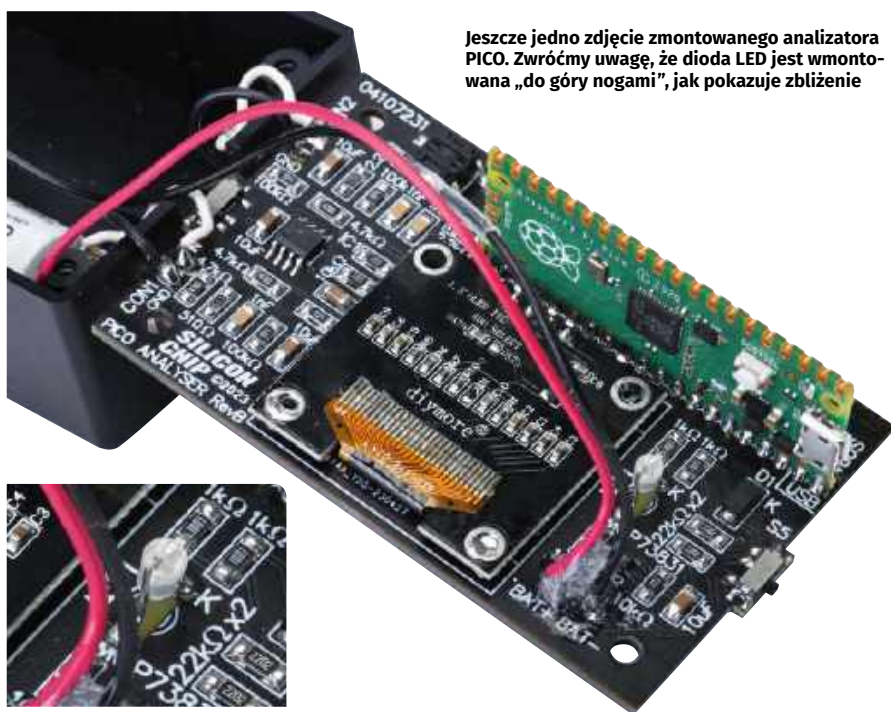
Analizator będzie w pełni gotowy do użycia po przeprowadzeniu prostej kalibracji. Bez kalibracji analizator też zadziała, ale jego dokładność będzie gorsza.

Będziemy potrzebowali multimetru lub oscyloskopu, którym będzie można dokładnie zmierzyć sygnał zmienny o wartości skutecznej 500 mVsk, oraz kabla z dwiema wtyczkami RCA do utworzenia pętli z wyjścia do wejścia. My zrobiliśmy krótki kabelek łączący jedynie środkowe zaciski wtyczek; połączenie masy jest w tym przypadku realizowane na płytce drukowanej.

Zasilamy analizator z kabla USB, co uruchomi ładowanie akumulatora. Przez kilka



Jeśli nie macie kabla RCA-RCA, możecie wykonać prosty kabel zwrotny lutując krótki przewód łączący środkowe piny dwóch wtyczek RCA. Taki kabel będzie niezbędny do testowania i kalibracji analizatora. Przydała się nam również para wtyczek RCA z przewodami połączeniowymi, umożliwiającą podłączenie analizatora do płytki prototypowej w celu eksperymentowania



Jeszcze jedno zdjęcie zmontowanego analizatora PICO. Zwróćmy uwagę, że dioda LED jest wmontowana „do góry nogami”, jak pokazuje zbliżenie

sekund będzie wyświetlany ekran powitalny. W tym czasie stabilizują się napięcia polaryzujące.

Naciskamy przycisk MODE („tryb”), aż pojawi się ekran SETTINGS („ustawienia”; **ekran 1**), a następnie wciskamy przycisk OK. **Ekran 2** pokazuje pierwszy krok kalibracji – INPUT OFFSET („offset wejścia”). Sprawdzamy, czy nic nie jest podłączone do wejścia i czekamy, aż wartość widoczna w czwartym wierszu ustabilizuje się na stałym poziomie. Naciskamy przycisk DOWN, a następnie OK.

Ekran 3 przedstawia kalibrację OUTPUT LEVEL („poziom wyjściowy”). Analizator wytworzy standardowy przebieg sinusoidalny 500 mVsk. Należy mierzyć napięcie tego przebiegu na wyjściu CON2. Przyciskami UP („w górę”) i DOWN („w dół”) zmieniamy współczynnik kalibracji, aż na mierniku odczytamy 500 mVsk. Naciskamy OK.

Na **ekranie 4** widać ustawianie INPUT LEVEL („poziom wejściowy”) na zakresie $\times 1$. Kabelkiem RCA-RCA dołączamy wyjście CON2 do wejścia CON1 i ustawiamy przełącznik S6 na pozycji $\times 1$. Analizator wie, że powinien odbierać sygnał 500 mVsk, więc łatwo może obliczyć współczynnik kalibracji. Po wyświetleniu komunikatu „DOWN to set” należy sprawdzić, czy w trakcie kalibracji był używany prawidłowy sygnał, a następnie nacisnąć DOWN, aby zatwierdzić obliczony współczynnik. Jeśli komunikat się nie pojawia, a S6 jest ustawiony prawidłowo, może to oznaczać jakiś problem z montażem płytki, np. użycie rezystora o niewłaściwej wartości. Ten krok kalibracji zależy również od prawidłowego ustawienia

amplitudy sygnału wyjściowego 500 mVsk w poprzednim kroku. Aby przejść do następnego ekranu, naciskamy OK.

Ekran 5 kalibruje INPUT LEVEL („poziom wejściowy”) na zakresie $\times 10$. Pozostawiamy zatem podłączony kabel i przełączamy S6 na zakres $\times 10$. Pojawi się komunikat jak w poprzednim kroku. Zatwierdzamy wynik kalibracji, wciskając DOWN.

Ekran 6 służy do zapisywania ustawionych parametrów w pamięci FLASH. W celu dokonania zapisu należy nacisnąć przycisk DOWN. Powinien zostać wyświetlony komunikat informujący o tym fakcie. Jeśli ustawienia zostały w jakiś sposób zniekształcone, można przywrócić ustawienia domyślne przyciskiem UP.

Działanie

Pozostałe ekrany pokazują tryby robocze. Do przełączania trybów służy przycisk MODE („tryb”), a przyciski UP („w górę”), DOWN („w dół”) i OK umożliwiają nawigowanie w ramach każdego trybu. Zmienianą wartość wyróżnia zazwyczaj para nawiasów kątowych $\langle \rangle$.

Podczas przełączania między trybami $\times 1$ i $\times 10$ należy również ręcznie zmienić tryb wejściowy na stronie SETTINGS („ustawienia”). Naciśnięcie przycisku DOWN powoduje wybranie trybu $\times 1$ (i użycie współczynnika kalibracji dla tego trybu), natomiast naciśnięcie UP wybiera tryb $\times 10$. Ostatni wiersz na tej stronie pokazuje bieżący współczynnik skalowania.

W prawym górnym rogu strony SETTINGS („ustawienia”) po włączeniu przełącznika

zasilania S5 wyświetlane jest napięcie akumulatora. Należy uważać, aby analizator nie pozostawał załączony, gdy nie jest używany, ponieważ wówczas nic nie zapobiega nadmiernemu rozładowaniu akumulatora.

Pierwszy tryb roboczy – WAVE OUTPUT („wyjście przebiegu”; **ekran 7**) – zadaje sygnał wyjściowy. Sygnał będzie zgodny z ostatnimi ustawieniami, chyba że kontrolę nad wyjściem będzie musiał przejąć inny tryb. Może się to zdarzyć gdy uruchamiany jest tryb SWEEP lub też gdy strona SETTINGS musi wygenerować swój przebieg kalibracyjny. Przycisk OK przełącza pomiędzy różnymi parametrami, a przyciski UP i DOWN zmieniają ich wartości.

Nastawiana częstotliwość może leżeć w zakresie od 10 Hz do 10 kHz. Przy niższych częstotliwościach krok ich ustawiania jest odpowiednio mniejszy. Moduł Pico zawiera generator kwarcowy, zatem nie przewidzieliśmy kalibracji częstotliwości. Dokładność częstotliwości kwarcu wynosi około 30 ppm (0,003%) i jest wystarczająca.

Poziom wyjściowy można ustawiać w krokach 50 mV (wartość międzyszczytowa lub skuteczną). W zależności od wybranego ustawienia wyświetlane są odpowiednie wartości. Stosunek wartości międzyszczytowych do skutecznych zmienia się w zależności od kształtu fali.

Sygnały wyjściowe o amplitudach do około 2 Vpp powinny być niezniekształcone, nie podlegając tym samym ograniczeniom wzmacniacza operacyjnego, zależnym od obciążenia wyjścia. Użyty wzmacniacz operacyjny jest dość odporny na błędy i wytrzymuje dowolnie długie zwarcie wyjścia.

Kolejny parametr wybiera przebieg sinusoidalny, prostokątny, trójkątny, piłokształtny lub biały szum. Parametr ostatni umożliwia wyłączenie lub załączenie sygnału bez zmiany innych ustawień. Gdy wyjście jest wyłączone, analizator nadal wytwarza przebieg, ale o amplitudzie 0 V.

W trybie SPECTRUM („widmo”) wyświetlane jest widmo przebiegu wejściowego (**ekran 8**). Przyciski UP i DOWN zmieniają skalę poziomą, przycisk OK przełącza pionową skalę amplitud między PEAK (wartości szczytowe) i TOTAL (wartości skuteczne). Do transformaty Fouriera stosowane jest okienkowanie, a zatem nawet czysta fala sinusoidalna będzie się zazwyczaj rozciągać na wiele przedziałów częstotliwości. Więcej informacji na temat okienkowania można znaleźć w Silicon Chip w artykule „Low Frequency Distortion Analyser” z kwietnia 2015 roku (www.siliconchip.au/Article/8441). Obliczane częstotliwości

prążków są interpolowane z wielu przedziałów i z powodu błędów zaokrąglenia mogą nieco odbiegać od faktycznych.

Tryb SCOPE („oscylloskop”; ekran 9) pokazuje kształt przebiegu wejściowego, tak jak na oscyloskopie. Przyciski UP i DOWN zmieniają skalowanie poziome (podstawę czasu), a przycisk OK przełącza między wykresem kropkowanym i liniowym. Gdy wyświetlanych jest kilka okresów fali, lepszą czytelność da tryb liniowy.

Skalowanie pionowe odbywa się automatycznie, w oparciu o zmierzoną amplitudę międzyszczytową, wyświetlaną po lewej stronie wykresu. „Oscyloskop” stara się wyzwać oscylogram przy dodatnim przejściu przebiegu przez zero, a jeśli wyzwalanie się nie powiedzie, wyświetlany jest po prostu ostatnio pobrany fragment sygnału.

Analiza harmoniczna

HARMONIC ANALYSIS („analiza harmoniczna”; ekran 10) dostarcza informacji o częstotliwości podstawowej sygnału, zawartości harmonicznych i zmierzonym współczynniku THD (całkowite zniekształcenia harmoniczne). W tym trybie przyciski UP, DOWN i OK nie są używane.

Jeśli zmierzmy sinusoidalny sygnał wyjściowy analizatora, otrzymamy wartość THD około 1%, z czego około 0,7% wprowadza stopień wyjściowy, a 0,3% – stopień wejściowy. Wartości te są zależne od częstotliwości.

Ostatnim trybem jest SWEEP („przemiatanie” częstotliwości i charakterystyka częstotliwościowa). Ekran 11 przedstawia stronę ustawień tego trybu, a ekran 12 – wyniki. Dolna i górna częstotliwość są potęgami liczby 10 i można je ustawiać w zakresie od 10 Hz do 10 kHz. Ilość punktów częstotliwości może wynosić maksymalnie 30. Pomiar w każdym punkcie trwa około 1/3 sekundy.

Jest opcja przełączania między „przemiataniem” pojedynczym a powtarzanym w pętli. Domyślna liczba punktów częstotliwości to 10. Tyle jest na wykresie widocznym na ekranie 12. Wyjście analizatora podłączono tu do jego wejścia.

Pozioma skala częstotliwości jest logarymiczna. Linie przerywane siatki odpowiadają drugiej i piątej dekadzie każdej dekady. Skala pionowa jest nastawiana przyciskami UP i DOWN. Na siatce dodatkową linią przerywaną jest oznaczony poziom –3 dB.

W ramach testowania podłączyliśmy między wejściem a wyjściem prosty obwód filtra dolnoprzepustowego RC (z rezystorem 1 kΩ i kondensatorem 1 μF). Zgodnie z oczekiwaniem, tryb SWEEP pokazał spadek charakterystyki do –3 dB przy częstotliwości



Ekran 1. Po pierwszym załączeniu analizatora naciskamy MODE aż przejdziemy do strony SETTINGS w celu wykonania kalibracji. Aby ją rozpocząć, naciskamy OK



Ekran 2. Gdy jest znajdowany poziom INPUT OFFSET, wejście musi pozostawać niepodłączone. Czekamy aż wyświetlana wartość się ustabilizuje. Wartość tę zapisujemy przyciskając DOWN, a następnie wciskamy OK



Ekran 3. Woltomierzem napięcia zmiennego lub innym przyrządem mierzymy wartość skuteczną sygnału na wyjściu i operujemy przyciskami UP i DOWN, aż miernik wskaże 500 mV. Następnie naciskamy OK



Ekran 4. Aby wykonać kolejne kroki kalibracji, łączymy kablem z wtykami RCA wyjście z wejściem. Upewniamy się, że przełącznik zakresu S6 jest ustawiony na 1x. Po wyświetleniu wiersza za-chęty naciskamy DOWN, a następnie OK



Ekran 5. Zgodnie z instrukcjami na ekranie, ustawiamy przełącznik na 10x. Jeśli S6 jest w nie-właściwej pozycji lub sygnał nie jest wykrywany, pojawi się komunikat. Współczynnik zatwierdza-my przyciskiem DOWN, a następnie OK



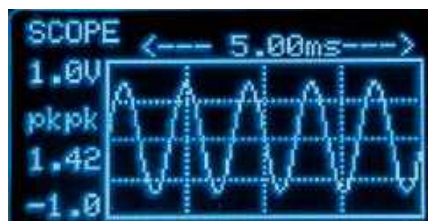
Ekran 6. Ekran ten pojawi się, gdy naciśniemy OK ponownie. Wciskamy DOWN, aby zapisać wartości parametrów kalibracji w pamięci FLASH. Pojawi się komunikat potwierdzający wykonanie tej czynności



Ekran 7. Na stronie WAVE OUTPUT naciśnięcie OK powoduje przechodzenie między parametrami, a przyciski UP i DOWN modyfikują je. Kształt przebiegu wyjściowego można ustawiać również poprzez port szeregowy USB



Ekran 8. Na stronie SPECTRUM przyciski UP i DOWN są używane do zmiany skalowania poziomego, natomiast OK przełącza skalę pionową między poziomem szczytowym i całkowitym



Ekran 9. Strona SCOPE również wykorzystuje przyciski UP i DOWN do zmiany skalowania poziomego. OK przełącza między wyświetlaniem punktowym i liniowym



Ekran 10. ANALIZA HARMONICZNA dostarcza informacji o zawartości harmonicznych w badanym przebiegu. Dobrym sposobem sprawdzenia tej funkcji jest dołączenie wejścia analizatora do wyjścia

Wady wewnętrznego przetwornika ADC

Nasz projekt analizatora PICO opierał się początkowo na dość optymistycznych założeniach. Mikrokontroler w module RP2040 ma 12-bitowy przetwornik analogowo-cyfrowy (ADC), mieliśmy więc nadzieję z użyciem nadpróbkowania uzyskać wyniki prawie jak dla rozdzielczości 14-bitowej. Jednak po podłączeniu do analizatora sygnału wyjściowego z Audio Precision System One (o współczynniku THD+N wynoszącym około 0,0004%), odczyt THD wynosił około 0,3%, czyli mniej więcej jak dla ośmiu bitów rozdzielczości.

Uważne przestudiowanie karty katalogowej RP2040 ujawniło erratę odnoszącą się do bloku ADC. Było tam stwierdzenie, że parametr ENOB (efektywna liczba bitów) przetwornika jest w rzeczywistości bliski ośmiu.

Przetwornik ADC opiera się na rejestrze kolejnych przybliżeń (SAR), a do pomiaru napięcia w układzie scalonym wykorzystane są niewielkie kondensatory o wartościach proporcjonalnych do kolejnych potęg dwójki. Całkowita pojemność kondensatorów wynosi około 1 pF, co oznacza, że najmniejsze z nich są rzędu femtofaradów (fF czyli 10⁻¹⁵ F)!

W toku dokładnych testów ustalono, że wartość niektórych z tych kondensatorów jest za niska o około 0,8%, poczynawszy od trzeciego najbardziej znaczącego bitu (o wadze 512; przypis redaktora). Patrz <https://pico-adc.markomo.me/INL-DNL/>. Pracownicy Raspberry Pi Foundation stwierdzili, że jest to spowodowane rozbieżnością między wynikami ich symulacji projektowych a rzeczywistością realizacją kondensatorów próbujących na kostce krzemowej.

Przetestowaliśmy efekt na naszym sprzęcie. Zmodyfikowaliśmy tymczasowo program, aby zliczał, ile razy każda wartość ADC (4096 możliwości) pojawiła się w zestawie próbek. Następnie wytworzyliśmy w analizatorze przebieg trójkątny. Przebieg taki przyjmuje teoretycznie każdy możliwy poziom trwający przez ten sam odcinek czasu, ponieważ nachylenie przebiegu w każdym półokresie jest stałe.

Wynik tej analizy – histogram – przedstawia **rysunek 4**. Zwróćmy uwagę na zliczenia zerowe na każdym końcu, odpowiadające wartościom, które w prawdziwym przebiegu nie występowały. Istnieją niewielkie szczyty w pobliżu skrajnych wartości przebiegu, ponieważ tam nachylenie sygnału się zmieniało (przebieg lekko się zaokrąślał).

Cztery wyraźne szczyty na wykresie – ogólnie dość płaskim – pokazują, że przetwornik ADC „sądzi”, iż przebieg spędza więcej czasu przy tych właśnie wartościach. Problematyczne wartości ADC to 511, 1535, 2559 i 3583. Wszystkie one wskazują na problemy z trzecim bitem MSB.

Wniosek jest taki, że przetwornik ADC niedokładnie mierzy napięcia w pobliżu tych punktów. Wejście może się zmieniać o około 10 kroków przetwornika, a wartość wyjściowa się nie zmienia. W rezultacie odczyt może być odchyłony od idealnego nawet o pięć kroków. Przetwornik nie reaguje liniowo.

W inny sposób ilustruje to wykres INL (nieliniowości całkowitej) z karty katalogowej RP2040 (**rysunek 5**). Wykres pokazuje odchylenia dokładności rzeczywistego przetwornika ADC od dokładności przetwornika idealnego. W praktyce linia powinna być dość płaska.

Oprogramowanie analizatora zostało uzupełnione o etap korekcji, w którym próbuje się skompensować nieliniowość przetwornika ADC. Sprowadza to mierzony współczynnik THD z 0,4% do 0,3%. Użyty sposób korekcji pokazano na **rysunku 6**.

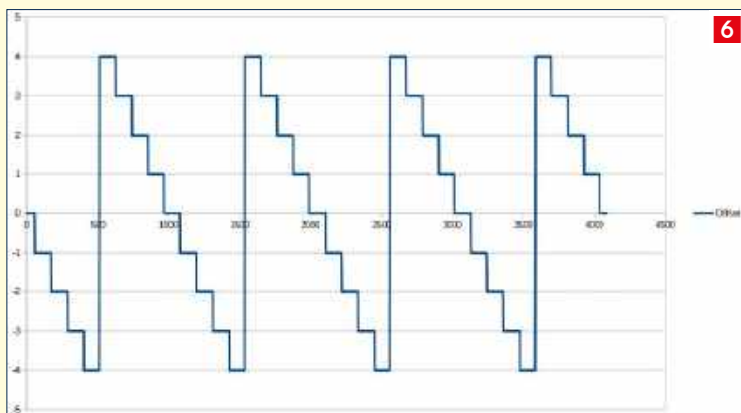
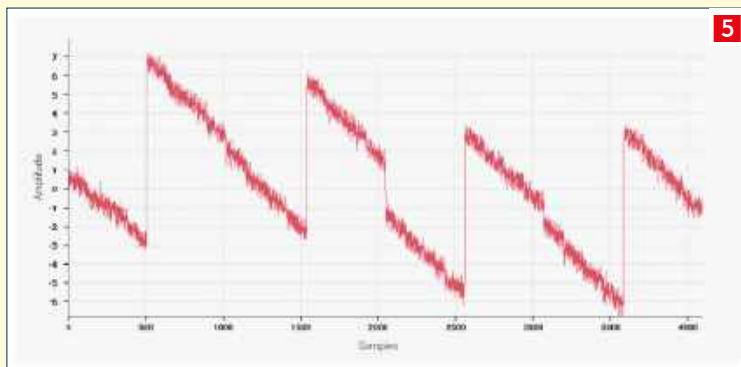
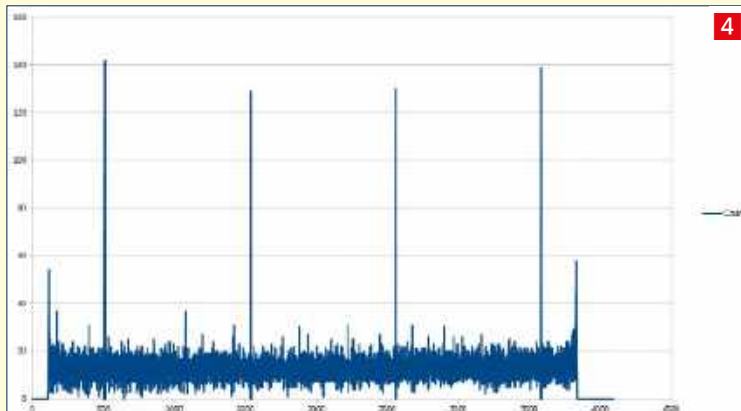
Korekcja sprawia, że wymienione wyżej cztery wartości ADC zostają niejako „porozpraszcane” w zakresie wartości przetwornika. Dzięki temu wykres jest ogólnie bardziej liniowy. Nadal jednak nie zmienia to faktu, że wartości te zajmują szerszy zakres napięć niż inne.

Rysunek 6 przypomina wykres INL. Próbowaliśmy jako korektę zastosować ten właśnie rysunek oraz kilka innych charakterystyk, w tym niektóre korygujące mniejsze błędy przy innych bitach przetwornika. W praktyce wybraliśmy ten z rysunku 6, ponieważ dawał najlepszą poprawę odczytów zniekształceń.

Wartości korekcji są przechowywane w tablicy o nazwie ADCADJ w pliku util.h. Aby zobaczyć efekt bez stosowania korekty, można w pliku zakomentować wywołania funkcji ADCfix().

Opisywany błąd wykazują wszystkie mikrokontrolery w modułach RP2040 znajdujących się obecnie w obiegu. Niewykluczone, że w przyszłych wersjach układów scalonych błąd ten zostanie poprawiony, a korekcja stanie się niepotrzebna.

Z tego wszystkiego płynie następujący wniosek: uważnie czytajmy karty katalogowe!





Ekran 11. W trybie SWEEP przyciski UP, DOWN i OK działają podobnie jak w trybie OUTPUT. Istnieje możliwość uruchomienia „przemiatania” pojedynczego lub wykonywania go w pętli



Ekran 12. Tutaj przyciski UP i DOWN zmieniają skalowanie pionowe. Nieoznakowana pozioma linia przerywana to poziom -3 dB w odniesieniu do poziomu wyjściowego analizatora

około 160 Hz i dalsze obniżanie charakterystyki przy częstotliwościach wyższych.

Przy bezpośrednim połączeniu wyjścia z wejściem odpowiedź częstotliwościowa powinna być idealnie płaska. Występują jednak niewielkie spadki przy 10 Hz i 10 kHz, bo zaczynają działać filtry dolno- i górnoprzepustowe. Mały „podskok” w okolicach 20 Hz jest efektem ubocznym funkcji okienkowania.

Naciśnięcie przycisku OK na stronie wykresu kończy pomiar w pętli. Jeśli wyświetlany jest komunikat <OK>, nastąpi powrót do menu ustawień.

Sterowanie z komputera

Na module Raspberry Pi Pico jest port USB. Użyliśmy go do realizacji dodatkowych funkcji sterowania oraz wyprowadzania wyników. Zalecamy korzystanie z programu

terminala, np. TeraTerm (pod Windows) czy minicom (w Linuksie), ponieważ monitor łączy szeregowego w środowisku Arduino jest trochę za prosty.

Polecenia wprowadzane z klawiatury komputera, symulujące funkcje uruchamiane przyciskami analizatora, działają natychmiastowo. Inne polecenia wymagają potwierdzenia klawiszem Enter.

By uzyskać pełną listę poleceń, należy wpisać „?” i nacisnąć Enter. Klawisze wymienione na dole ramki (ekran 14) symulują cztery przyciski na płycie głównej.

Większość poleceń symuluje elementy sterujące trybu WAVE OUTPUT („wyjście przebiegu”). Działają one również wtedy, gdy aktywny jest inny tryb, mogą zatem zaoszczędzić czas potrzebny na przełączanie się między trybami w celu zmiany ustawień.

	A	B	C	D
1	Sweep results:			
2	Index	Frequency (Hz)	Amplitude (V)	Gain (dB)
3	0	10	0.293	-4
4	1	21	0.55	0
5	2	46	0.497	0
6	3	100	0.499	0
7	4	215	0.499	0
8	5	464	0.498	0
9	6	999	0.498	0
10	7	2154	0.494	0
11	8	4641	0.474	0
12	9	10000	0.397	-2

Ekran 13. Polecenie „d” wysłane z terminala szeregowego uruchamia zrzut wyników w formacie CSV. Tu widzimy wyniki z trybu SWEEP, wklejone do arkusza kalkulacyjnego.

```

Commands (followed by Enter):
A: set RMS amplitude in mV (eg a1000)
P: set peak-to-peak amplitude in mV (eg p2000)
F: set frequency in Hz (eg f440)
W1: output sinewave
W2: output triangle wave
W3: output square wave
W4: output sawtooth wave
W5: output white noise
O0: output off
O1: output on
D: dump data from Scope, Spectrum, Harmonic or Sweep modes
< and > emulate DOWN and UP
/ and M emulate OK and MODE

```

Ekran 14. Lista poleceń dostępnych z wirtualnego portu szeregowego USB, wyświetlona w programie TeraTerm. Listę można wyświetlić poleceniem „?”

Wartość skuteczną sygnału wyjściowego (w miliwoltach) ustawia polecenie rozpoczynające się od „a” lub „A”, po którym następuje liczba. Na przykład „a500” zadaje poziom wyjściowy 500 mVsk. Analogicznie przez „p” lub „P” ustawiamy wartość międzyszczytową (w miliwoltach). Opcja „f”/„F” zadaje częstotliwość w hercach, polecenie „w”/„W” ustawia kształt fali, a „o”/„O” wyłącza lub załącza przebieg na wyjściu.

Należy pamiętać, że ustawienie zbyt wysokich wartości parametrów może skutkować zniekształceniem przebiegów.

Kolejne polecenie – „d”/„D” – realizuje w trybach SCOPE, SPECTRUM, HARMONIC ANALYSIS i SWEEP „zrzucenie” do komputera wartości próbek z najbliższego pobrania ich do bufora. Dane są w formacie CSV (wartości oddzielone przecinkami), dzięki czemu można je bezpośrednio wklejać do arkuszy kalkulacyjnych obsługujących format CSV. W przypadku trybu SWEEP zrzut nastąpi po zakończeniu następnego „przemiatania”. Na ekranie 13 widzimy dane z ekranu 12, wklejone do arkusza kalkulacyjnego.

I wreszcie polecenie „~”. Powoduje ono zresetowanie modułu Pico. Podczas wydawania tego polecenia należy trzymać wciśnięty przycisk BOOTSEL. Spowoduje to przejście do trybu bootloadera, umożliwiające przeprogramowanie modułu.

Wnioski

Analizator PICO to proste, zwarte urządzenie, które poza samym modulem Raspberry Pi Pico wykorzystuje niewiele innych elementów. Parametry analizatora nie są wygórowane. Sądzymy jednak, że prostota i niski koszt tego przyrządu sprawią, że okaże się narzędziem przydatnym. ■

Tim Blythman



Materiały dodatkowe dostępne są na stronie Silicon Chip:
<https://www.siliconchip.com.au/Shop/8/6771>

Materiały dodatkowe są również dostępne na stronie elportal.pl/do-pobrania

Artykuł reprodukowano na podstawie umowy z magazynem „Silicon Chip”, 2022.
www.siliconchip.com.au

262 144 sposoby na „Grę w życie”

Od dawna fascynuje mnie „Gra w życie” (Game of Life), wymyślona przez Johna Conwaya. Zapraszam do udziału w niezwykłej przygodzie, podczas której przekształcimy tę łamigłówkę z lat 70. w ekscytujący spektakl wizualny, wykorzystując matrycę RGB LED i specjalne metody programowania. Będziemy mogli przeprowadzać symulacje, stosując każdy sensowny zestaw reguł!

Opisywany projekt to historia składająca się z wielu elementów. Niektóre krążyły w moim umyśle przez 50 lat, łącząc się w końcu w bardzo satysfakcjonującą całość. Zaczęło się od „Gry w życie”, o której po raz pierwszy usłyszałem pod koniec lat 70. Brałem wtedy udział w moim pierwszym (i jedynym) kursie programowania – w języku BASIC, przy użyciu terminala DECwriter podłączonego do minikomputera PDP-11 w mojej szkole średniej. Jednym z ćwiczeń było napisanie programu do „Gry w życie”. Był to przykład stosowania w BASIC-u zagnieżdżonych pętli FOR...NEXT – świetne, zabawne ćwiczenie dla początkującego programisty, w wyniku którego powstawały fascynujące zmieniające się wzory. W kolejnych latach nie mijała moja fascynacja tą grą. A gra przenosiła się kolejno z terminali minikomputerów na monitory PC-tów, laptopy i iPady.

Gra w życie

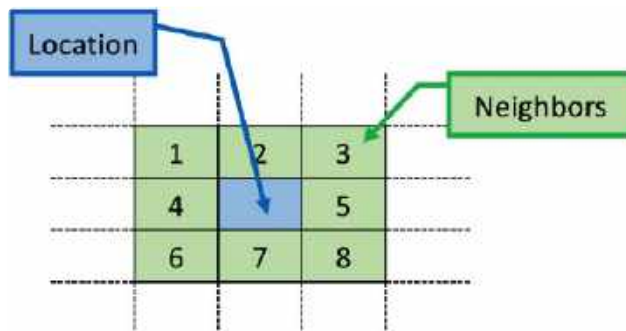
W skrócie: „Gra w życie” [1] jest prostą symulacją automatu komórkowego, opartą na określonych regułach. Rozgrywana jest na prostokątnej planszy o dowolnym rozmiarze. Elementami gry są komórki, rozlokowane w kratkach planszy. W trakcie gry, w kolejnych pokoleniach, komórki mogą trwać, umierać lub się rodzić. Przeżycie lub wyginiecie danej komórki w następnym pokoleniu zależy tylko od tego, ile sąsiadów ma ta komórka w pokoleniu bieżącym. Na prostokątnej planszy każda kratka ma od zera do ośmiu sąsiadów (rysunek 1).

Oryginalne zasady sformułowane przez Conwaya są następujące:

- komórka przetrwa do następnego pokolenia, jeśli ma dwóch lub trzech sąsiadów; komórki z mniejszą liczbą sąsiadów „umierają z samotności”, komórki z większą liczbą sąsiadów „umierają z przełudnienia”;

Wykaz elementów:

Adafruit Matrix Portal – wyświetlacz internetowy sterowany przez CircuitPython (Adafruit 4745);
matryca LED 64×32 RGB – odstęp 5 mm (Adafruit 2277);
3×3-cyfrowy przetwornik przyciskowy z kodowaniem BCD (np. Littlefuse 3P-3-23-3-1-0-2);
ekspander wejścia/wyjścia MCP23017;
4 rezystory 2 kΩ;
przełącznik przyciskowy;
przełącznik jednobiegunowy (SPST);
zasilacz USB;
obudowa z plexi.



Rysunek 1. Kratka na planszy i jej ośmiu sąsiadów

- w pustej kratce urodzi się w następnym pokoleniu nowa komórka, jeśli kratka ta ma jako sąsiadów dokładnie trzy komórki.
- Conway wybrał te zasady w toku eksperymentów, tak aby były spełnione następujące kryteria [2]:
- nie powinno być takiego układu początkowego, dla którego istnieje prosty dowód na to, że populacja może rosnąć bez ograniczeń;
 - powinny istnieć układy początkowe, które najwyraźniej rosną bez ograniczeń;
 - powinny istnieć proste układy początkowe, które rosną i zmieniają się przez dłuższy czas, zanim zakończą się na trzy możliwe sposoby:
 - całkowicie zanikając – z powodu przełudnienia lub samotności,
 - przechodząc w układ stabilny, który od tej pory pozostaje niezmienny,
 - wchodząc w fazę oscylacji, w której układy tworzą niekończący się cykl o długości dwóch lub więcej pokoleń.

W 1970 roku gra była rozgrywana pionkami na szachownicy zwykłych warcabów. Wkrótce jednak cały proces został skomputeryzowany, a innowacje trwają do dziś [3].

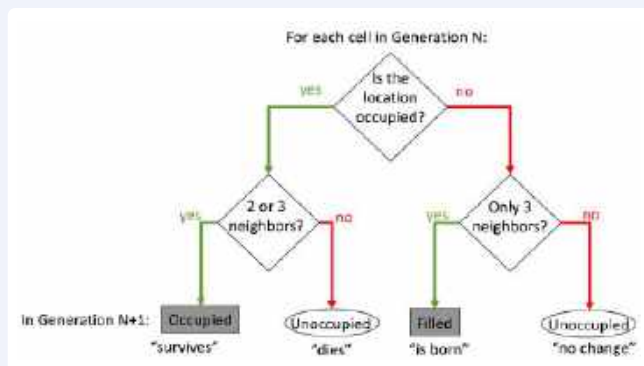
Ciekawe algorytmy

Kolejnym elementem historii „Gry w życie” było ukazanie się książki „Nowy rodzaj nauki” („A New Kind of Science”) Stephena Wolframa, w której omówił on różne algorytmy automatów komórkowych, tworzących wzory często bardzo interesujące wizualnie. Jedną z kluczowych części analizy Wolframa było zakodowanie reguł przejścia do następnego pokolenia w postaci liczb dwójkowych. Kodując reguły w ten sposób, możliwe jest zbadanie konsekwencji wszystkich możliwych zestawów reguł w dobrze zdefiniowany i powtarzalny sposób [4]. Skłoniło mnie to do zastanowienia się nad sposobem kodowania reguł „Gry w życie”, tak aby możliwe było zbadanie wszystkich możliwych zestawów tych reguł. Przekonałem się, że reguły można zakodować w 18-bitowej liczbie dwójkowej (szczegóły w ramce tekstowej *Kodowanie reguł na sześciu cyfrach ósemkowych*). **Przypis redaktora:** przypominamy, że zapis liczby binarnej w systemie ósemkowym polega na podziale jej zapisu na grupy 3-bitowe oraz zastąpieniu każdej takiej grupy odpowiednią cyfrą od 0 do 7. Każdy z dziewięciu najmłodszych bitów (od 0 do 8) określa stan pustej kratki w następnym pokoleniu (1 = zajęta przez komórkę; 0 = pusta), mającej od zera do ośmiu sąsiadów – bit 0 określa stan w następnej generacji, jeśli kratka ma obecnie zero

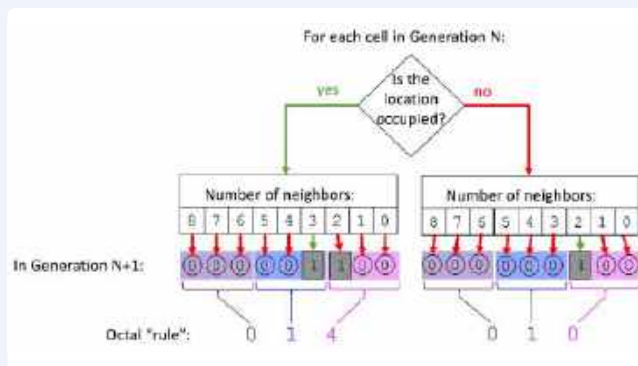
Kodowanie reguł na sześciu cyfrach ósemkowych

Oryginalne zasady „Gry w życie” można opisać w skrócie:

- komórka przeżywa, jeśli ma 2 lub 3 sąsiadów, w przeciwnym razie umiera;
- komórka rodzi się w każdej pustej kratce, która ma dokładnie trzech sąsiadów.



Te same reguły można również przedstawić w następujący sposób, gdzie „1” w regule oznacza kratkę z komórką w następnym pokoleniu (przeżycie lub narodziny) dla danej liczby sąsiadów, a „0” w regule oznacza kratkę pustą w następnym pokoleniu (śmierć lub pozostanie pustą).



Wzorec bitowy określa stan komórki w następnym pokoleniu; indeks we wzorcu bitowym (numer bitu) to liczba sąsiadów w połączeniu z tym, czy w danej kratce jest obecnie komórka, czy nie. Poniżej jako przykład przedstawiono dekodowanie reguły 2560208:

Tabela 1. Interpretacja zestawu reguł ósemkowych „2560208”

Tabela 11. Interpretacja zestawu reguł ósemkowych „200020																		
kratka	z komórką									pusta								
sąsiedzi	8	7	6	5	4	3	2	1	0	8	7	6	5	4	3	2	1	0
następne pokolenie	0	1	0	1	0	1	1	1	0	0	0	0	0	1	0	0	0	0
kod ósemkowy	2			5			6			0			2			0		

W tłumaczeniu na tekst byłoby to:

żywe komórki przeżywają, jeśli mają 1, 2, 3, 5 lub 7 sąsiadów; w przeciwnym razie umierają;

komórka rodzi się w każdej pustej kratce, która ma dokładnie 4 sąsiadów.

Korzystając z tego formatu kodowania można określić dowolny zestaw reguł dla „Gry w życie” z 262144 możliwych.

sąsiadów, bit 1 – gdy ma jednego sąsiada itd. Analogicznie, dziewięć najstarszych bitów (od 9 do 17) określa stan w następnym pokoleniu kratki zajętej przez komórkę – bit 9 przy zerze sąsiadów; bit 10 przy jednym itd. Korzystając z tego sposobu kodowania możliwe jest określenie dowolnego zestawu reguł opartych na liczbie sąsiadów danej kratki. Co ciekawe, reguły te nie muszą mieć „biologicznego sensu”. Na przykład zestaw reguł 2560208, w którym komórki utrzymują się, jeśli mają 1, 2, 3, 5 lub 7 sąsiadów. Ten zestaw reguł daje interesujące zachowanie, mimo że nie ma biologicznego powodu, dla którego 4, 6 i 8 sąsiadów powodowałoby przepełnienie, natomiast 3, 5 i 7 nie.

Ponieważ wszystkie możliwe zestawy reguł można zakodować na 18 bitach, a $2^{18} = 262144$, to znaczy, że istnieje 262144 sposobów „Gry w życie”. Można zbudować urządzenie, które pozwoliłoby użytkownikowi badać każdą z tych możliwości i obserwować, jak populacja każdego z różnych „gatunków” rozwijałaby się w czasie.

Ostatni element projektu powstał w wyniku wielu rozmów z artystką elektroniczną Kelly Heaton [5], którą poznałem dzięki artykułowi w Elektor Magazine [6]. Kelly zachęciła mnie do wyjścia poza kanon i do tworzenia dzieł, które są zarówno wierne algorytmowi jak i atrakcyjne wizualnie. Doprowadziło mnie to do odejścia od standardowego schematu kolorów „Gry w życie”, w którym komórki są kolorowe, a pola puste – czarne. Chciałem uwidocznić charakter zachodzących zmian. Pokolorowałem więc komórki nowo narodzone na niebiesko; komórki, które utrzymują się przez więcej

niż jedno pokolenie, są zielone, natomiast komórki, które umierają, pozostawiają fioletowego „ducha” na jedno pokolenie („duchy” nie są liczone jako sąsiedzi). Zostało to szczegółowo przedstawione na **rysunku 2**.

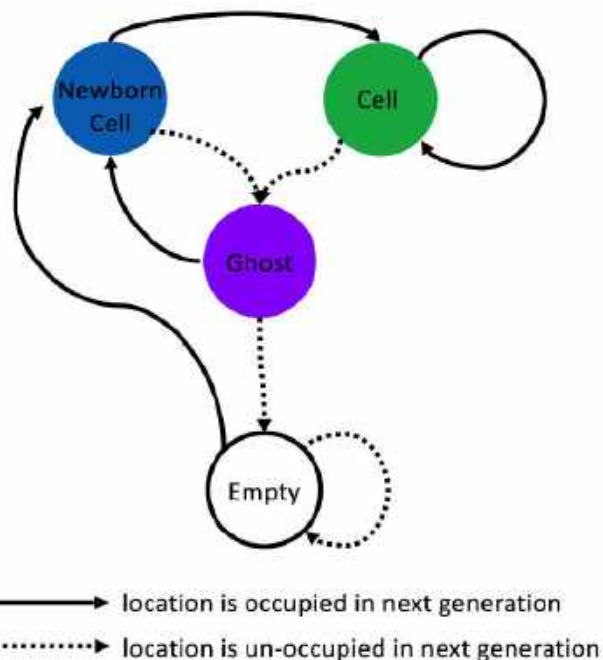
Sprzęt

Końcowy element projektu – sprzęt – przybył jako prezent świąteczny od moich synów: sterownik Adafruit Matrix Portal [7] i matryca LED RGB 32×64 [8]. Całość zmontowałem w obudowie z plexi, aby były widoczne wszystkie elementy wewnętrzne. Końcowy efekt pokazano na **rysunku 3**. Po lewej stronie znajdują się trzy zestawy po trzy przełączniki kodowe, którymi ustawia się reguły symulacji i inne parametry pracy. Górny i środkowy zestaw przełączników nastawia reguły w postaci sześciu cyfr ósemkowych – każda cyfra ósemkowa to 3 bity, a 6 cyfr daje wymagane 18 bitów. Trzy górne cyfry określają reguły dla kratek zawierających komórki, a więc reguły określające, które komórki przetrwają do następnego pokolenia. Środkowe trzy cyfry określają reguły dla pustych krater, a więc reguły narodzin nowych komórek. Przykładowo – oryginalne reguły Conwaya zostałyby wyrażone jako 014 – 010. Ramka tekstowa Kodowanie reguł na sześciu cyfrach ósemkowych opisuje to bardziej szczegółowo. Pod wyświetlaczem znajduje się – od lewej do prawej – płytka ekspandera wejścia/wyjścia MCP23017, płytka sterownika Matrix Portal, przycisk resetowania i przełącznik trybu wyświetlania.

Każde uruchomienie rozpoczyna się od losowego rozmieszczenia komórek na planszy. Następnie odbywa się symulacja przez ustaloną liczbę pokoleń, po czym następuje ponowne uruchomienie z nowym losowym układem komórek. Szczegóły tego procesu określają trzy dolne cyfry przełączników. Lewa cyfra określa gęstość organizmów w pierwszym pokoleniu; konkretnie – szansa na zajęcie na początku każdej z krtek wynosi 10% razy ustawienie tego pokrętki. Tak więc ustawienie „3” oznacza, że świat będzie na początku wypełniony komórkami w około 30%. Niektóre zestawy reguł dają bardziej interesujące wyniki z większą albo mniejszą liczbą komórek. Środkowa cyfra określa liczbę pokoleń, zanim odbędzie się ponowne uruchomienie. Wybrałem skalę wykładniczą, co umożliwia szeroki zakres nastawień. Konkretnie – symulacja działa przez $2 \cdot 10^{(N/2)}$ pokoleń; zatem np. przy ustawieniu 4 będzie 200 pokoleń. Niektóre wzory znikają szybko, podczas gdy inne nie znikają nigdy. Ustawienie ilości pokoleń pomaga dobrać optymalny czas symulacji. Wreszcie – cyfra prawa określa w sekundach opóźnienie między kolejnymi pokoleniami. Ustawienie 0 daje maksymalną prędkość. Spowolnienie umożliwia szczegółowe śledzenie wpływu reguł na przebieg symulacji. Ostatnie dwa elementy sterujące to przycisk, który natychmiast ponownie losuje układ komórek i uruchamia symulację oraz przełącznik do wyboru albo kolorów standardowych (czarny = puste pole, niebieski = komórka) albo wielokolorowego schematu mojego autorstwa.

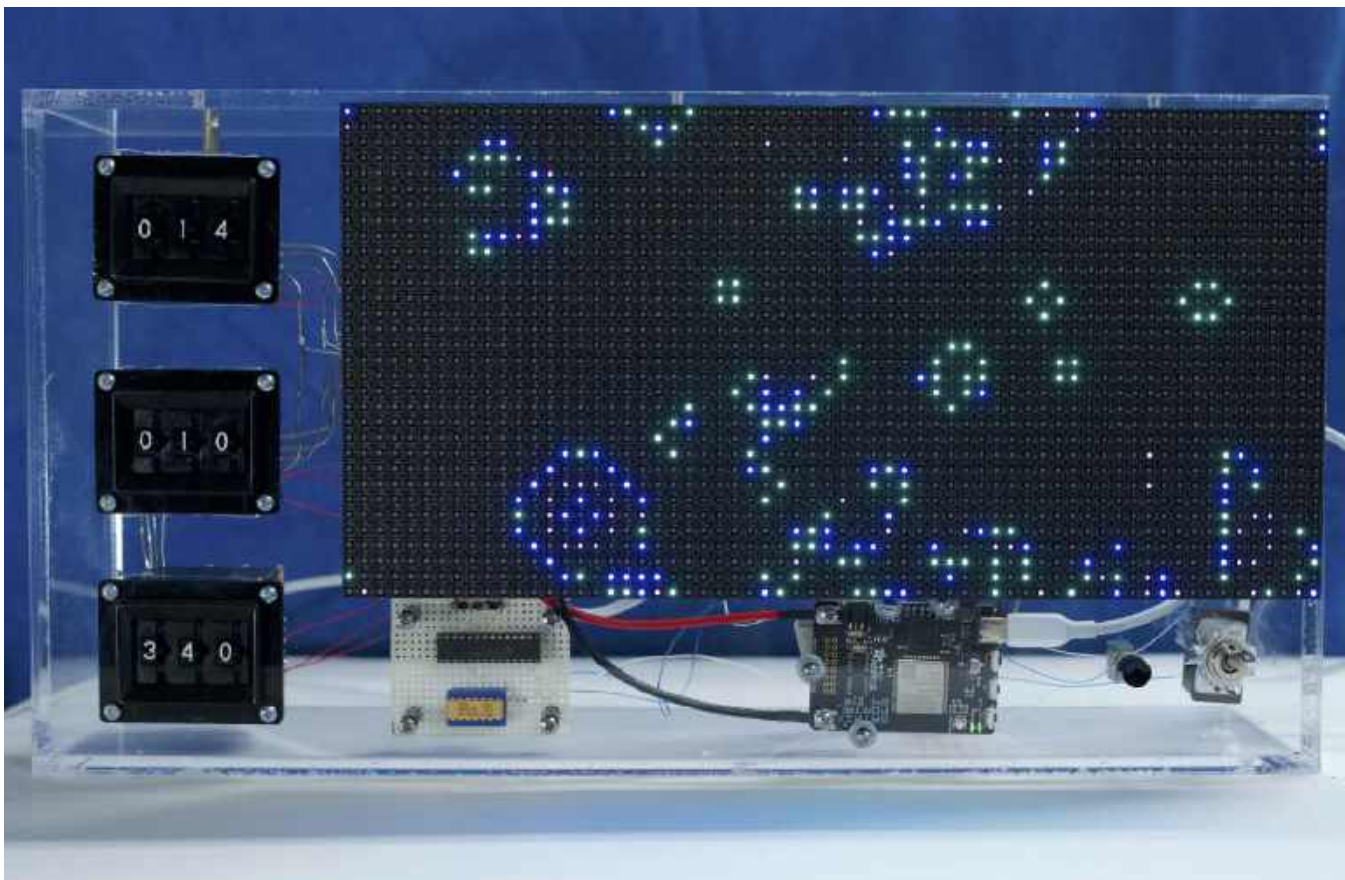
Zrób to sam

Schemat na **rysunku 4** jest bardzo prosty. Matrix Portal i wyświetlacz są podłączone zgodnie z instrukcjami Adafruit dla sterownika i matrycy LED. Pokazane na rysunku elementy dodatkowe są zasilane z +3,3 V i komunikują się przez I²C z Matrix Portal za pośrednictwem



Rysunek 2. Diagram stanów dla trybu wielokolorowego wyświetlania komórek

złącza Stemma QT. Sporo zadań jest wykonywanych przez ekspander wejścia/wyjścia I²C (MCP23017). Wyjścia tego układu załączają kolejno przełączniki BCD (SW1...SW3, SW6...SW11), a wyjścia BCD



Rysunek 3. Demonstracja gotowego urządzenia

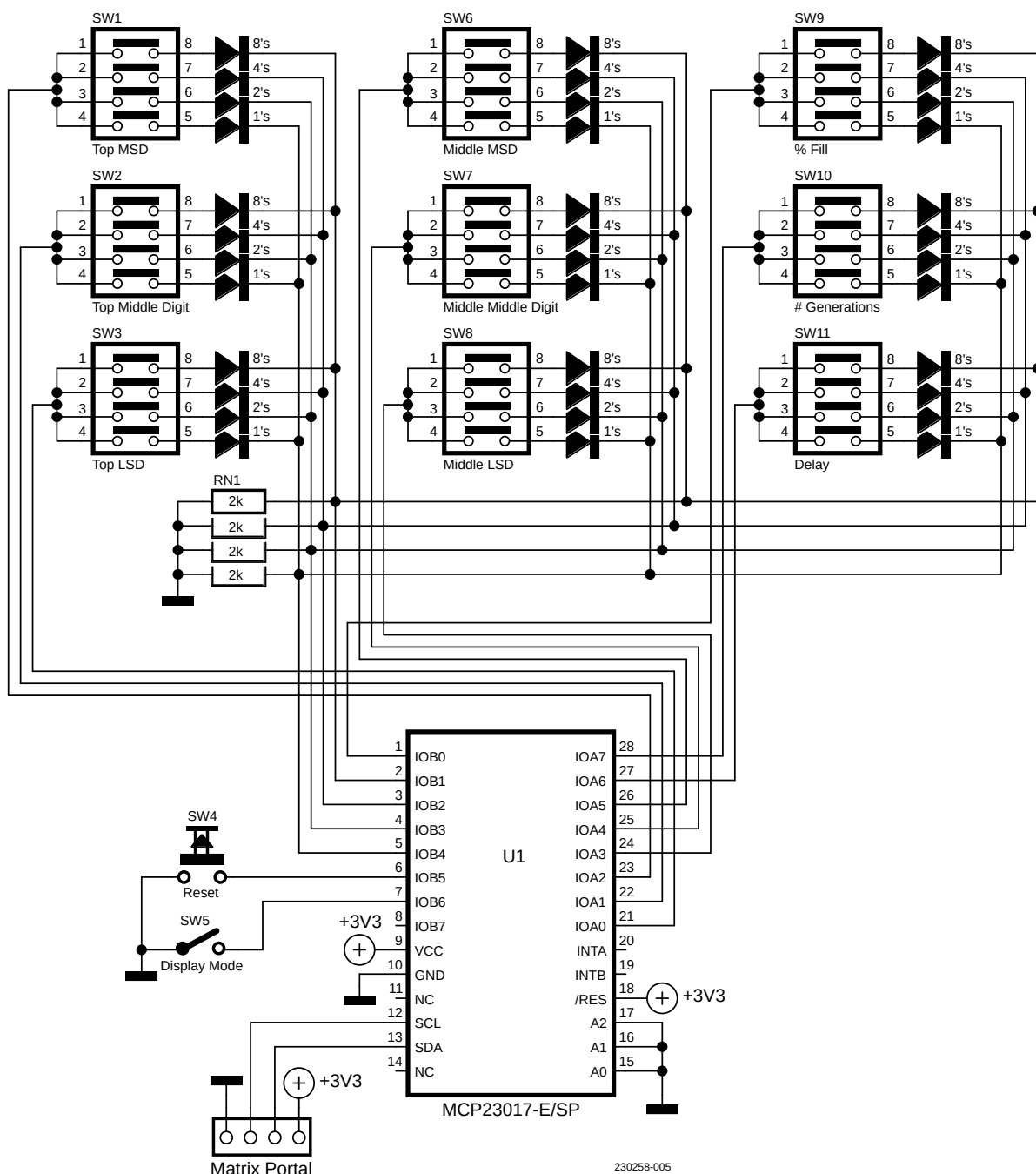
załączonego przełącznika są następnie odczytywane przez wejścia ekspandera. Ponieważ MCP23017 nie ma rezystorów ściągających do masy, na liniach BCD dołączono rezystory zewnętrzne. Dwa wejścia MCP23017, z aktywnym wewnętrznym podciąganiem do +3,3 V, są używane do odczytu przełączników resetowania i trybu wyświetlania. Zasilanie jest dostarczane przez zasilacz 5 V/4 A, który zasila Matrix Portal przez kabel USB-C.

Program jest oparty na kodzie „Gry w życie” dostarczonym przez Adafruit dla Matrix Portal and Display [9], który, aby pętla programu działała bardzo szybko, wykorzystuje kilka bardzo sprytnych sztuczek. Program zmodyfikowałem, aby obowiązywały reguły oparte na liczbach ustawionych przełącznikami zamiast reguł „standardowych”. Krótko mówiąc – wczytuję stany przełączników i używam ich

do utworzenia 18-bitowej listy jedynek i zer reprezentujących stan kratki w następnym pokoleniu, w zależności od aktualnej zajętości kratki i liczby sąsiadów („duchy” nie są uwzględniane jako sąsiedzi). Biorę liczbę sąsiadów, dodaję 9 jeśli w kratce jest komórka, i używam tej sumy jako indeksu do listy reguł. Stosuję również wzorec wielokolorowy, opisany wcześniej. Program jest dostępny na stronie projektu Elektor Labs [10].

Idąc dalej

Wyniki są fascynujące. Stworzyłem playlistę na YouTube [11] z kolekcją krótkich filmów pokazujących wzory, które pojawiają się przy różnych ustawieniach przełączników. Niezwykle jest to, że niektóre zmiany w regułach mają duży wpływ na sposób, w jaki wzory rosną,



Rysunek 4. Schemat – nie obejmuje Matrix Portal, wyświetlacza 64×32 i zasilania USB

zmieniają się, poruszają i wymierają, podczas gdy inne dają jedynie subtelne zmiany. Dopiero zacząłem badać możliwości (262144 to dość duża liczba!). Mam na razie kilka spostrzeżeń:

- ustawienie bitów najmłodszych (0 i 9) – co umożliwia narodziny i przetrwanie przy braku sąsiadów – prowadzi do radykalnych zmian (odwróceń) dodatnich i ujemnych co pokolenie;
- ustawienie bitów starszych – kontrolowanie urodzeń i przeżycia z większą liczbą sąsiadów – ma zwykle mniejszy wpływ, być może dlatego, że komórki z wieloma sąsiadami występują rzadko;
- oryginalny zestaw reguł – 0140108 – zwykle ewoluuje w struktury stosunkowo stabilne lub figury oscylujące; inne reguły, takie jak 1141108, są „ulubieńcami” mojej żony, ponieważ wydają się nigdy nie rozwiązywać;
- można tworzyć reguły, w których struktury są tworzone i stopniowo wypełniane w interesujący sposób, jeśli zezwoli się na przetrwanie przy dużej ilości sąsiadów, a narodziny tylko przy niewielu; na przykład 2560208.

Urządzenie wisi w tej chwili przy naszym stole kuchennym i często okazuje się, że ktoś w domu znalazł nowy zestaw reguł, który daje fascynujące wzory.

W miarę jak pracujemy nad tym modelem i udostępniamy go różnym osobom, zyskaliśmy kilka pomysłów dla tych, którzy chcieliby zbudować swój własny egzemplarz. Po pierwsze, interesujące byłoby dodanie do wielokolorowego schematu większej liczby kolorów. Na przykład, kolory mogłyby stopniowo zmieniać się poprzez więcej odcieni w miarę „dojrzwania” komórek. Kelly Heaton zasugerowała również, by kolory przechodziły jeden w drugi płynnie, a nie zmieniały się skokowo. Wreszcie, moi domownicy nie znający systemu ósemkowego sugerują, aby znaleźć sposób na bardziej przejrzyste pokazanie w regułach zależności między bitami a liczbą sąsiadów. Można by dodać rząd 18 LED-ów i dwójkowo wyświetlać odpowiednik ustawień ósemkowych. A jeszcze prościej – użyć do ustawiania reguł 18 indywidualnych przełączników, po jednym dla każdego bitu.

Mam nadzieję, że Czytelnicy tego artykułu zainspirują się do zbudowania egzemplarza tego urządzenia – albo z dedykowanym wyświetlaczem i sterownikiem na mikrokontrolerze, albo całkowicie na komputerze – i podzielą się zestawami reguł, które uznają za interesujące.

Kto by pomyślał, że dzielenie się 6-cyfrowymi liczbami ósemkowymi może być tak interesujące?... ■

Brian White (Stany Zjednoczone)

Linki:

[1] Wikipedia: „Gra w życie” Hortona Conwaya: https://en.wikipedia.org/wiki/Conway%27s_Game_of_Life

REKLAMA

O Autorze

Brian White jest profesorem na Wydziale Biologii Uniwersytetu Massachusetts w Bostonie. Ukończył studia biologiczne na MIT i Stanford. Jako nauczyciel prowadzi również badania w dziedzinie biologii i edukacji biologicznej oraz opracowuje powiązane oprogramowanie. Brian jest również entuzjastą elektroniki, muzykiem i radioamatorem (KA1TBQ). Więcej szczegółów na temat jego projektów elektronicznych można znaleźć na stronie [12].



- [2] Martin Gardner, „MATHEMATICAL GAMES: The fantastic combinations of John Conway’s new solitaire game ‘life’,” Scientific American 223 (October 1970): 120-123: <https://web.stanford.edu/class/sts145/Library/life.pdf>
- [3] LifeWiki: <https://conwaylife.com/wiki>
- [4] Stephen Wolfram, A New Kind of Science; zobacz zwłaszcza rozdział: <https://wolframscience.com/nks/p53--more-cellular-automata>
- [5] Kelly Heaton Studio: <https://kellyheatonstudio.com>
- [6] C.J. Abate, „Making Art with Electricity: A Q&A With Kelly Heaton”, Elektor Circuit Special 2022: <https://elektormagazine.com/articles/making-art-with-electricity-a-q-a-with-kelly-heaton>
- [7] Matrix Portal – CircuitPython-Powered Internet Display: <https://adafruit.com/product/4745>
- [8] Matryca LED RGB 64×32 – rozstaw 5 mm: <https://adafruit.com/product/2277>
- [9] Przykład Adafruit: Conway’s „Game of Life”: <https://learn.adafruit.com/rgb-led-matrices-matrix-panels-with-circuitpython/example-conways-game-of-life>
- [10] Strona projektu na Elektor Labs: <https://elektormagazine.com/labs/262144-ways-to-play-the-game-of-life>
- [11] Filmy z Game of Life: https://youtube.com/playlist?list=PL2wCLIp1MkRpBHeZ2svMs-CdLvf2B_9N
- [12] Strona internetowa Autora: <https://brianwhite94.wixsite.com/electronics>

Pytania lub komentarze?

W przypadku pytań technicznych lub komentarzy dotyczących tego artykułu prosimy o kontakt z Autorem pod adresem brian.white@umb.edu lub z redakcją Elektor pod adresem editor@elektor.com.

Publikujemy dla projektantów i programistów elektroniki

ELPORTAL.pl

TAWOIA Glass (szkło kwarcowe)

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



BESTSELLERY sklepu AVT – sklep.avt.pl

3 unikalne
serie
gniazdek
i
włączników

Rabat dla Czytelników EdW
przy zakupie podaj kod **EdW2505GW**

Kod ważny do 30.09.2025

Rabat dla Prenumeratorów EdW
przy zakupie podaj numer prenumeraty

-5%

-10%

Ceramic Loft (ceramika)

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



Retro PRL (bakelit)

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



Generator sygnału testowego FM na pasmo 2 m

Opisywany generator testowy wytwarza sygnały w paśmie 144...148 MHz (2 m), wykorzystując w nieco nietypowy sposób układ DDS AD9834 taktowany częstotliwością znacznie niższą, bo 80 MHz. Generator może wytwarzać sygnał CW (fala ciągła) o trzech różnych poziomach, z opcjonalną modulacją częstotliwości.

Układ scalony bezpośredniej syntezy cyfrowej (DDS) Analog Devices AD9834 jest powszechnie stosowany w generatorach funkcyjnych i oscylatorach w.c.z. Jego wariant układowy opisywany w niniejszym artykule wytwarza programowane cyfrowo sygnały FM na kanałach pasma amatorskiego 2 m. Układ może być zasilany nawet z pojedynczego ogniwa AA.

Pasmo amatorskie 2 m obejmuje częstotliwości z zakresu 144...146 MHz w Europie i 144...148 MHz w Ameryce Północnej, Australii i Nowej Zelandii.

Projekty oparte na niedrogim układzie AD9834 są publikowane od ponad 20 lat. Gdybyśmy jednak przyjrżeli im się bliżej, to stwierdzilibyśmy, że prawie wszystkie z nich naśladują rozwiązania z not aplikacyjnych Analog Devices, a różnice polegają głównie na odmiennym typie mikrokontrolera sterującego, wyświetlacza czy innym układzie płytki drukowanej.

Kilka lat temu zbudowałem na AD9834 prostą płytkę testową z mikrokontrolerem ATtiny45 i paroma przyciskami. Sprawdziłem kilka układów filtrów wyjściowych, przetestowałem także możliwości „podkręcenia” częstotliwości taktującej AD9834. Szybko ustaliłem, że układy AD9834-BRUZ mogą być taktowane zegarem do 85 MHz, co stanowi istotny wzrost w stosunku do standardowych 50 MHz podawanych w karcie katalogowej. Nieco droższe podtypy układu z końcówką oznaczenia -CRUZ są specyfikowane na 75 MHz i podobno pracują aż do 100 MHz.

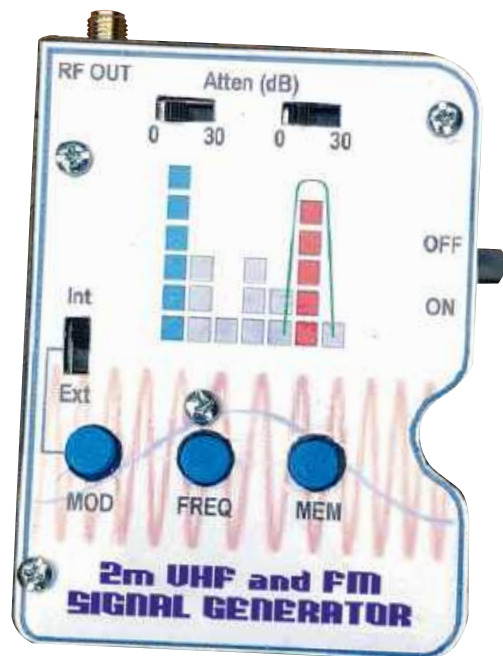
Na tych eksperymentach skończyły się wówczas moje działania z układem AD9834. Nie miałem dla niego natychmiastowego zastosowania i odłożyłem go na półkę. Spodziewałem się zatem różnych niespodzianek, gdy niedawno znajomy poprosił mnie o potwierdzenie wyników, jakie zmierzył w swoim generatorze sygnałów z AD9834.

Modulacja fazy

Gdy jeszcze działałem z układem AD9834, moją uwagę zwróciły funkcje modulacji częstotliwości (FM) i modulacji fazy (PM), wymienione w karcie katalogowej. W innych moich projektach generatorów DDS używałem układów AD9850 i AD9851. Ich rejestry fazy mają ograniczone możliwości, co sprawia, że układy te zasadniczo nie nadają się do zaimplementowania modulacji FM.

Oparte na AD9834 projekty generatorów o zmiennej częstotliwości (VFO), z jakimi się stykałem, wykorzystywały układ tylko jako strojony oscylator w.c.z. lub generator funkcyjny na zakres akustyczny. Nigdy wcześniej nie widziałem żadnej wzmianki o modulacji fazy w AD9834. Badając sprawę bliżej przekonałem się, że AD9834 ma dla celów modulacji fazy specjalne rejestry 12-bitowe, co znacznie zwiększa możliwości modulacji PM/FM niż w przypadku układów AD9850/9851. Rozbudziło to moją ciekawość.

W ciągu godziny zaimplementowałem w AD9834 modulację fazy, działającą jako tako w ograniczonym zakresie. Z karty katalogowej nie wynikało jednak jasno, jak uzyskać



określony poziom modulacji. Po wyczerpujących przeszukiwaniach aplikacji Analog Devices okazało się, że aplikacje te w dużej mierze milczą na ten temat.

Gdy badałem temat modulacji fazy, niezbędnych wyliczeń i innych funkcji układu, zdałem sobie sprawę, że można osiągnąć efekt zupełnie inny niż można by oczekiwać od zwykłego projektu na AD9834. Rezultatem jest niniejszy nowatorski generator testowy CW/FM na pasmo 2 m.

Częstotliwości DDS i aliasy

AD9834 to układ generatora DDS typowy dla szerokiej klasy układów produkowanych przez Analog Devices. Schemat blokowy układu jest przedstawiony na **rysunku 1**. Układ zawiera dwa programowane rejestry częstotliwości i dwa rejestry fazy, co umożliwia szybki wybór jednej z dwóch kombinacji częstotliwości i fazy.

Rejestry te sterują zaawansowanym generatorem sterowanym cyfrowo (NCO), opartym na 28-bitowym akumulatorze fazy. Wyjście NCO steruje tablicą funkcji sinus i 10-bitowym przetwornikiem cyfrowo-analogowym, w którym odbywa się konwersja słów wyjściowych generatora NCO na analogowy przebieg sinusoidalny.

Specyfikacja

- Częstotliwość wyjściowa 144...148 MHz, ustawiana w krokach 500 kHz
- Pamięć na cztery dodatkowe częstotliwości pasma 2 m, programowane przez użytkownika
- Poziom wyjściowy -45 dBm, -75 dBm i -105 dBm
- Fala nośna bez modulacji lub z modulacją FM
- Dewiacja FM ± 3 kHz (ton 1 kHz)
- Wejścia akustyczne modulacji zewnętrznej i dla kodera CTCSS
- Analogowy filtr wyjściowy, łatwy do zestrojenia; pozostałe funkcje są cyfrowe

Częstotliwość wyjściową generatora można obliczyć ze wzoru $f_{OUT} = NREG \cdot f_{CLK} / 228$. f_{CLK} jest częstotliwością zewnętrznego zegara taktującego, a $NREG$ to 28-bitowa wartość cyfrowa wpisana do jednego z dwóch przełączanych rejestrów częstotliwości. Rejestry fazy umożliwiają przesuwanie fazy sygnału wyjściowego o programowany kąt fazowy.

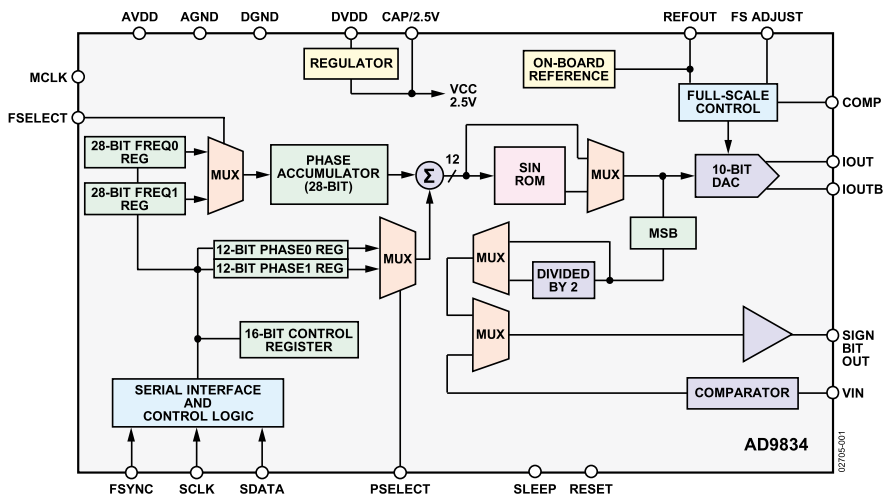
Przy użyciu zewnętrznego kwarcowego generatora zegarowego i z odpowiednim filtrem dolnoprzepustowym (LPF), przebieg wyjściowy AD9834 jest całkiem czystą i niskoszumną sinusoidą. Jej częstotliwość może wynosić do około 30 MHz. Nastawianie częstotliwości jest możliwe z rozdzielczością 0,3 Hz. Dokładność częstotliwości zależy od precyzji i stabilności generatora zegarowego.

Sygnał wyjściowy wytwarzany przez układ DDS jest w rzeczywistości znacznie bardziej złożony. W przypadku braku filtra na wyjściu, DDS obok sygnału podstawowego generuje również szereg składowych dodatkowych. Mogą one rozciągać się znacznie poza 300 MHz, co pokazano na **rysunku 2** (rysunek z karty katalogowej AD9834, wzbogacony kolorami).

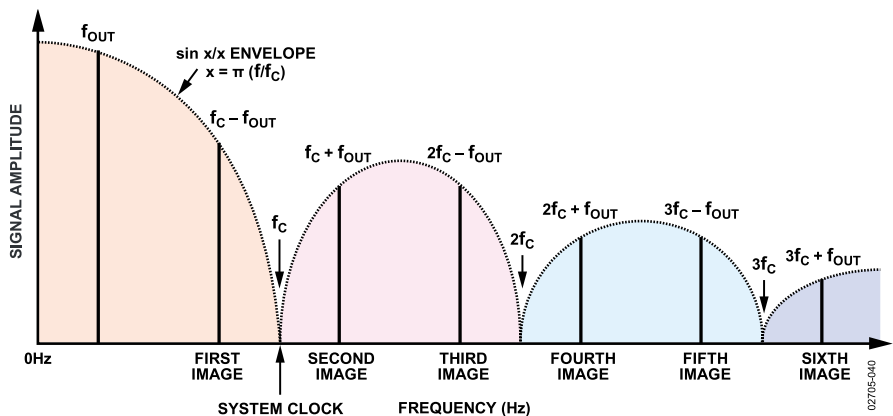
Filtr LPF przepuszcza zwykle bez przeszkód podstawowy sygnał wyjściowy (f_{OUT}), natomiast silnie tłumi inne, niepożądane sygnały. Częstotliwość wyjściowa może być ustawiona na dowolną wartość aż do połowy częstotliwości zewnętrznego zegara; na przykład 40 MHz, gdy DDS jest taktowany zegarem 80 MHz.

Niepożądane składowe wytwarzane przez AD9834 nazywane są „aliasami” ([ang. „image”; przypis redaktora](#)). Pierwszy z „aliasów” ma częstotliwość $f_{CLK} - f_{OUT}$. Jeśli zegar taktujący DDS wynosi 80 MHz, a f_{OUT} jest zmieniana (poprzez zmianę wartości w rejestrze częstotliwości) w krokach od 1 Hz do, powiedzmy, 30 MHz, ten pierwszy „alias”, oznaczony na rysunku 2 jako $f_C - f_{OUT}$, zaczyna się od częstotliwości 80 MHz i z każdym krokiem jego częstotliwość maleje, aż osiąga 50 MHz ($= 80 \text{ MHz} - 30 \text{ MHz}$). Widać to bardziej szczegółowo na **rysunku 3**.

Jeśli częstotliwość pożądanego sygnału wyjściowego zbliża się do 40 MHz (czyli połowy częstotliwości zegara DDS), ten pierwszy „alias” obecny na wyjściu staje się coraz bardziej dokuczliwy. On również zbliża się do 40 MHz – od góry. Powinien zostać w jakiś sposób odfiltrowany. Staje się to coraz trudniejsze, gdy pożądaný sygnał wyjściowy przekracza 30 MHz i zbliża się do 40 MHz. W przypadku większości generatorów DDS wymaga to użycia filtra



Rysunek 1. Schemat blokowy AD9834 – kopia z karty katalogowej, uzupełniona kolorami. Jest to typowy generator DDS z 28-bitowym akumulatorem fazy. Przy zewnętrznym zegarze 80 MHz generator może wytwarzać czyste przebiegi sinusoidalne do częstotliwości 30 MHz

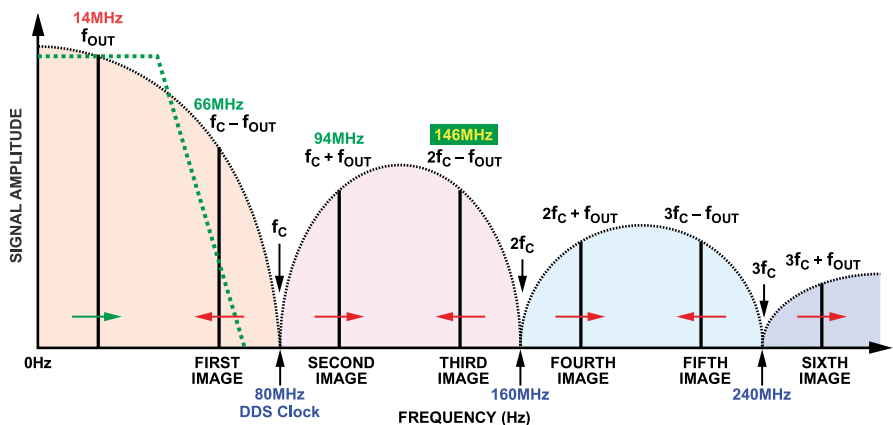


Rysunek 2. AD9834 bez żadnych dodatkowych filtrów na wyjściu wytwarza częstotliwości znacznie przekraczające 300 MHz

dolnoprzepustowego z bardzo stromą charakterystyką odcięcia. Taki filtr powinien zwykle zaczynać tłumić sygnał zaraz powyżej 30 MHz, a dalej tłumienie to ma gwałtownie wzrastać, osiągając przy 40 MHz wartość

co najmniej 60 dB. Dlatego większość konstrukcji zawiera na wyjściu złożony filtr dolnoprzepustowy na przykład 5. lub 7. rzędu.

Na rysunku 3 widać również, że poziom wyjściowy pożądanego sygnału DDS nie



Rysunek 3. Widmo sygnału wyjściowego AD9834, taktowanego zegarem DDS 80 MHz i z rejestrem częstotliwości ustawionym na 14 MHz. Strzałki pokazują kierunek przemieszczania się niefiltrowanych „aliasów” wraz ze wzrostem częstotliwości podstawowej 14 MHz. Zielona linia przerywana pokazuje charakterystykę typowego filtra dolnoprzepustowego wysokiego rzędu, zwykle używanego do usuwania niepożądanych składowych z sygnału wyjściowego

jest taki sam przy każdej częstotliwości. Poziom ten spada zgodnie z wartościami funkcji $\sin(x)/x$. Na przykład z zegarem 80 MHz poziom wyjściowy przy częstotliwości 30 MHz jest o 2 dB niższy niż przy 1 MHz. Gdy sygnał pożądaný ma częstotliwość 30 MHz, to poziom „aliasu” 50 MHz jest od niego słabszy tylko o 4,5 dB. Bez dobrego filtra sygnał wyjściowy AD9834 zawierałby przy 30 MHz silne zniekształcenia, spowodowane mieszaniem się sygnału pożądanego z pobliskim „aliasem” 50 MHz.

Praca AD9834 w paśmie 2 m (144–148 MHz)

Na rysunku 3 widzimy inne „aliasy”, powstające powyżej częstotliwości sygnału pożądanego i pierwszego „aliasu”. Jeśli sygnał pożądaný wynosi 14 MHz, a pierwszy „alias” 66 MHz, to „alias” następny wynosi 94 MHz. Kolejne „aliasy” są przy częstotliwościach 146 MHz, 174 MHz, 226 MHz i tak dalej.

Na wyjściu układu AD9834 pojawia się również śladowa liczba taktów zegara o częstotliwości 80 MHz. Przebieg zegarowy to fala prostokątna, więc ma również silną składową w postaci trzeciej harmonicznej; w tym przypadku przy 240 MHz.

Na rysunku 3 pokazano również kierunek przemieszczania się „aliasów” wraz ze wzrostem częstotliwości podstawowej generatora.

„Alias” mogą leżeć w radiowym paśmie 2 m. Jeśli na przykład częstotliwość podstawową AD9834 ustawimy kolejno na 16 MHz, 15 MHz, 14 MHz, 13 MHz i 12 MHz, a zegar taktujący DDS będzie wynosić 80 MHz, to trzeci „alias” wypadnie przy odpowiednio 144 MHz, 145 MHz, 146 MHz, 147 MHz i 148 MHz – aczkolwiek będzie miał znacznie niższy poziom niż częstotliwość podstawowa. Jeśli na wyjściu AD9834 damy

filtr pasmowy LC o stosunkowo wąskim paśmie przepustowym, dostrojony do zakresu 2 m, to możliwe będzie uzyskanie na wyjściu za filtrem wyłącznie sygnałów leżących w paśmie 2 m. Przykład charakterystyki takiego filtra, przeznaczonego do stosowania z AD9834, pokazano na **rysunku 4**.

Taki filtr wyjściowy musi być dopasowany do impedancji wyjściowej 200 Ω układu AD9834 i umożliwiać obciążenie wyjścia impedancją 50 Ω , typową dla zastosowań w.c.z.

Modulacja

Następnym krokiem – po zastosowaniu na wyjściu AD9834 filtra wybierającego częstotliwości pasma 2 m – było sprawdzenie, czy w układzie możliwe jest uzyskanie modulacji częstotliwości. Jak już wspominaliśmy, karta katalogowa AD9834 podkreśla możliwość modulacji fazy (a tym samym częstotliwości), ale nie podaje żadnych szczegółów. Również noty aplikacyjne Analog Devices nie zawierają żadnych szczegółów na temat tego, w jaki sposób rejestry fazy AD9834 mogą być wykorzystane do przeprowadzenia modulacji. Do tego, pomimo dokładnych poszukiwań, nie udało mi się znaleźć żadnego projektu generatora DDS, w którym funkcja modulacji zostałaby faktycznie zrealizowana.

Skloniło mnie to do samodzielnego zagłębienia się w temat modulacji fazy. Przeanalizowałem i przetestowałem działanie rejestrów fazy układu AD9834 w celu zrozumienia ich wpływu na sygnał wyjściowy.

Kiedy w latach 1970–1980 w paśmie amatorskim 2 m zaczęto stosować modulację częstotliwości (FM), wystąpiła tendencja do zastępowania jej modulacją fazy (PM). Pojawiły się twierdzenia o „lepszej jakości”

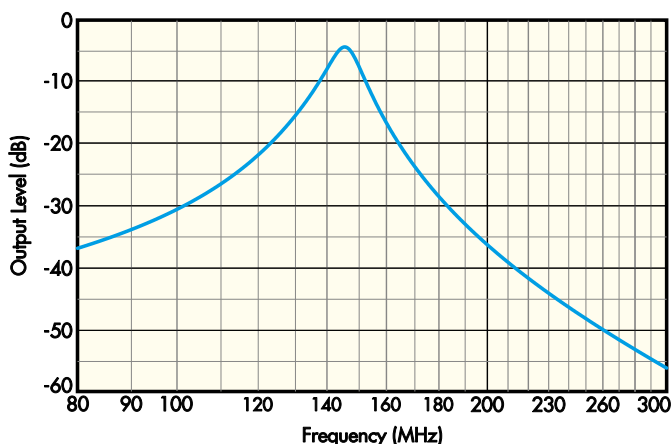
tego rodzaju modulacji i „bardziej naturalnie brzmiącej mowie”. Na poparcie tych twierdzeń było jednak niewiele dowodów.

Kwestie praktyczne, a szczególnie pojawienie się tanich diod pojemnościowych (warikapów), doprowadziły do wyparcia modulacji PM przez FM. Warikapy łatwo było zastosować w układach oscylatorów w.c.z., co redukowało liczbę użytych elementów, choć niestety czasami kosztem liniowości modulacji.

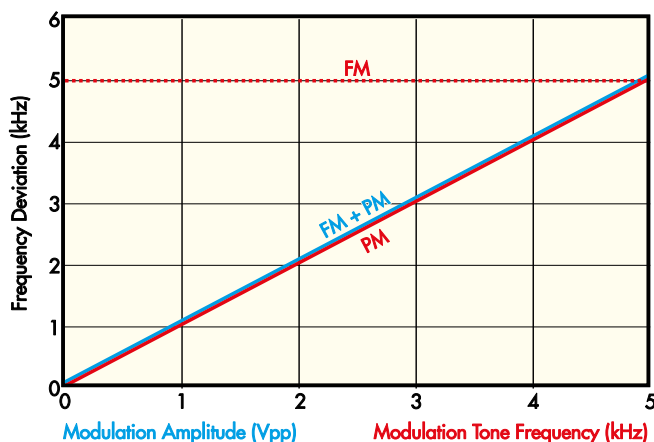
Modulacja PM szybko wypadła z obiegu, czego konsekwencją był spadek informacji o tym rodzaju modulacji w czasopiśmie technicznych i podręcznikach z tamtego okresu. Jednym z wyjątków był klasyczny „Radio Handbook” („Podręcznik techniki radiowej”) Williama Orra. Wydanie z 1981 roku zawierało zwięzły opis metody PM wraz z przykładami. Patrz podsumowanie na **rysunku 5**.

Przede wszystkim, dewiacja (stopień modulacji) częstotliwości nadajnika jest wprost proporcjonalna do amplitudy modulującego sygnału akustycznego, zarówno dla modulacji fazy, jak i częstotliwości. W przypadku FM dewiacja jest niezależna od częstotliwości sygnału modulującego. Natomiast w modulacji PM dewiacja wzrasta wraz ze wzrostem częstotliwości tego sygnału.

Stosowane początkowo nadajniki z modulacją fazową wykazywały specyficzną charakterystykę częstotliwościową dla sygnału modulującego, określaną później jako „preemfaza”. Jej istnienie wymagało zastosowania w odbiorniku dla sygnału zdemodulowanego odwrotnej charakterystyki częstotliwościowej, czyli „deemfazy”. Osiągało się to zwykle poprzez prosty obwód filtra dolnoprzepustowego RC bezpośrednio za stopniem dyskriminatora (detektora) odbiornika FM.



Rysunek 4. Filtr może mieć charakterystykę pozwalającą przepuszczać „aliasy” 144...148 MHz, a tłumić sygnał podstawowy, inne „aliasy” i zegar taktujący. Taką charakterystykę zapewni filtr pasmowoprzepustowy na dyskretnych cewkach i kondensatorach



Rysunek 5. Dewiacja częstotliwości w nadajniku z modulacją fazy lub częstotliwości zależy od poziomu modulacji. W nadajniku z modulacją fazy dewiacja zależy również od częstotliwości sygnału modulującego. Wykres adaptowany z podręcznika „Radio Handbook” Billa Orra, 1981 r.

Nadajniki FM, aby poprawnie współpracować z odbiornikami FM z „deemfazą”, wymagały dodania układu realizującego „preemfazę”. Zwykle tuż przed modulatorem FM dodawano obwód RC filtra górnoprzepustowego. W ten sposób niejako naśladowano modulację fazy.

Korzystnym efektem zrealizowania „preemfazy” w nadajniku FM (lub naturalne jej występowanie w nadajniku PM), a „deemfazy” w odbiorniku, jest zmniejszenie szumów, ponieważ ich słyszalność wzrasta z częstotliwością.

Modulacja fazy w AD9834

Wróćmy teraz do układu AD9834. Modulacja fazy jest w nim realizowana poprzez cykliczne zmiany wartości rejestru fazy (PHASE0/1 REG; rysunek 1). Wartość w tym rejestrze precyzyjnie przesuwają w fazie bieżący sygnał wyjściowy DDS.

Częstotliwość wyjściowa AD9834 (f_{OUT}) jest określona przez wartość rejestru częstotliwości. Przesunięcie fazy wnoszone przez 12-bitowy rejestr fazy wynosi $2\pi/4096$ radianów pomnożone przez wartość zawartą w tym rejestrze. Karta katalogowa nie podaje tego precyzyjnie, ale tak właśnie jest. Nie ma żadnych przykładów czy szczegółów zastosowania tego sposobu przesuwania fazy. Użytkownicy AD9834 muszą radzić sobie sami.

Jak się okazuje, poprzez cykliczne wpisywanie do rejestru fazy stanów proporcjonalnych do wartości sygnału modulującego możliwe jest wytworzenie pożądanego sygnału PM (a tym samym FM).

W tym momencie powinny zabrzmieć fanfary... tylko że jest jeszcze jeden istotny szczegół. Przesunięcie fazowe spowodowane modulacją fazy w AD9834 jest takie samo zarówno dla częstotliwości podstawowej, jak i dla wszystkich aliasów z rysunku 3. AD9834 całkowicie różni się pod tym względem od tradycyjnych modulatorów fazy i nadajników opisywanych w podręcznikach, od wczesnych nadajników radiofonii komercyjnej FM i od bardzo starych (przepraszam: „klasycznych”) nadajników amatorskich na pasmo 2 m. Zanim wprowadzono pętlę synchronizacji fazowej (PLL), tradycyjne nadajniki wykorzystywały do wytworzenia fali nośnej szereg stopni powielania (mnożenia) częstotliwości. Typowy wczesny nadajnik na pasmo 2 m miał generator kwarcowy 12 MHz, po którym następował ciąg stopni powielających częstotliwość. **Przypis redaktora: powielanie częstotliwości polegało na zniekształceniu wejściowej fali sinusoidalnej, czyli uzyskaniu częstotliwości harmonicznych, „wylawianiu”**

Szczegóły modulacji fazy w AD9834

Dewiacja częstotliwości nośnej przy modulacji fazy jest proporcjonalna zarówno do częstotliwości, jak i amplitudy sygnału modulującego, można zatem zastosować następujące równanie:

$$\text{Dewiacja częstotliwości [kHz]} = \text{przesunięcie fazowe [w radianach]} \cdot \text{częstotliwość modulacji [kHz]}$$

Jeśli na przykład częstotliwość sygnału modulującego wynosi 1 kHz i mamy przesunięcie fazy nośnej o $+0,5$ radiana, to wynikowa dewiacja częstotliwości sygnału wyjściowego wyniesie $+500$ Hz (pamiętamy, że 2π radianów = 360°).

Generator sygnału do testowania amatorskich odbiorników radiowych FM na pasmo 2 metrów (odstęp międzykanałowy 25 kHz) wykorzystuje standardowo ton testowy 1 kHz i dewiację częstotliwości nośnej ± 3 kHz. Stąd maksymalne przesunięcie fazowe fali nośnej w AD9834 wynosi $3 \text{ kHz} / 1 \text{ kHz} = 3$ [radiany].

12-bitowy rejestr fazy w układzie AD9834 wytwarza przesunięcie fazowe nośnej na wyjściu o π radianów, gdy PHASEREG = 2048. Aby uzyskać 3 radiany odchylenia fazy, do rejestru tego musi być wpisana wartość $3/\pi \cdot 2048 = 1956$.

Wewnętrzny generator w mikrokontrolerze wytwarza falę sinusoidalną o częstotliwości 1 kHz i wartości międzyszczytowej 3,7 Vpp, podawaną następnie do wejścia analogowego ADC0 mikrokontrolera. Wartość szczytowa w 10-bitowym rejestrze wyniku przetwarzania analogowo-cyfrowego ATtiny45 wynosi około 750. Napięcie odniesienia ADC to 5 V, zatem sygnał na wejściu ADC o amplitudzie 5 Vpp dałby wynik maksymalny, czyli 1023.

Oprogramowanie skaluje wynik przetwarzania ADC tak, aby z wartości 750 uzyskać wartość dla rejestru fazowego wynoszącą około 2250. Jest ona nieco wyższa niż wyliczona wartość 1956, a to z powodu błędów zaokrąglania w stosowanym prostym algorytmie obliczeń na liczbach całkowitych. Testy przy użyciu profesjonalnego miernika modulacji potwierdziły, że uzyskana wartość wywołuje w sygnale wyjściowym AD9834 dewiację ± 3 kHz.

W przypadku korzystania z zewnętrznego wejścia modulacji, maksymalna dewiacja częstotliwości, którą można osiągnąć, wynosi około $\pm 4,5$ kHz. Wynika to z ograniczenia pomiaru 10-bitowego przetwornika ADC wynoszącego 1023 (gdy na wejściu mamy amplitudę 5 Vpp). Tego poziomu wejściowego nie wolno przekraczać, w przeciwnym razie może dojść do uszkodzenia układu ATtiny45.

określonej (na przykład drugiej lub trzeciej) harmonicznej w obwodzie rezonansowym LC i doprowadzeniu jej do wyjścia stopnia powielającego. Zwykle używano jednego stopnia potrającego częstotliwość, po którym następowały dwa stopnie podwajaczy częstotliwości. Częstotliwość generatora 12 MHz była więc mnożona łącznie przez 12, co dawało końcowe 144 MHz. Również dewiacja częstotliwości 144 MHz była dwanaście razy większa niż częstotliwości 12 MHz. Była to korzystna cecha, ponieważ starsze modulatory fazy i częstotliwości nie zapewniały dużych wartości dewiacji. Ponieważ jednak stopnie mnożące częstotliwość w nadajniku mnożyły również dewiację fazy i częstotliwości, z łatwością osiągnąć na wyjściu wymagany poziom modulacji.

Inaczej wygląda sytuacja w przypadku generatora DDS. Jeśli częstotliwość wyjściowa AD9834 (f_{OUT}) jest ustawiona na 14 MHz, odpowiednia zmiana wartości rejestru fazy AD9834 spowoduje wytworzenie na wyjściu sygnału zmodulowanego tonem 1 kHz z dewiacją częstotliwości 3 kHz. **Przypis redaktora: w technice określenie „ton” dotyczy sygnału sinusoidalnego o częstotliwości akustycznej.** Ze względu na własności aliasingu powstającego w generatorze DDS, „alias” o częstotliwości 146 MHz będzie

miał identyczną dewiację częstotliwości: 3 kHz. Nie będzie „efektu mnożenia”, jaki występował w tradycyjnych nadajnikach PM i FM. Wszystkie częstotliwości „aliasów” są bezpośrednio i jednocześnie zależne od częstotliwości wyznaczonej w procesie syntezy DDS.

Modulacja fazy w AD9834 jest procesem w 100% cyfrowym, przeprowadzanym w generatorze. Maksymalna dewiacja częstotliwości jest ograniczona przez rozmiar rejestru fazy i proces syntezy.

W ramce „Szczegóły modulacji fazy w AD9834” opisano, w jaki sposób wartości w rejestrze fazy tworzą pożądaną dewiację PM (i FM).

Szczegóły układu

Na **rysunku 6** przedstawiono schemat miniaturowego generatora testowego FM na pasmo 2 m. Układ AD9834 jest sterowany z 8-pinowego mikrokontrolera ATtiny45 firmy Atmel/Microchip przez trójliniową magistralę szeregową SPI. ATtiny45 zawiera sprzętowy interfejs SPI, co upraszcza pisanie oprogramowania i pozwala na większą szybkość przesyłania danych. W układzie generatora interfejs SPI wykorzystuje tylko dwie linie (USCK i DO) z trzech, ponieważ nie jest konieczne odczytywanie danych z AD9834.



stotliwości sygnału wytwarzanego przez generator DDS została objaśniona w ramce „Szczegóły modulacji fazy w AD9834”. Dewiacja FM ± 3 kHz daje zwykle 60% szczytowego poziomu modulacji używanego do testowania kanałów o szerokości 25 kHz.

Sygnał modulujący jest do mikrokontrolera doprowadzany z zewnątrz. Dzięki temu można użyć zewnętrznego źródła modulacji. Przełącznik S4 należy w tym celu przełączyć w pozycję „EXT”. Maksymalny poziom wejściowy sygnału modulującego wynosi 5 Vpp, co odpowiada dewiacji $\pm 4,5$ kHz. Dla kanałów FM o szerokości 12,5 kHz odpowiedni będzie maksymalny poziom sygnału 3 Vpp, czyli dewiacja $\pm 2,5$ kHz, a poziom 2 Vpp wywoła dewiację $\pm 1,75$ kHz (60% odchylenia szczytowego). Wszystkie te wartości obowiązują przy założeniu, że napięcie zasilania wynosi 5 V.

Użyta w projekcie częstotliwość próbkowania 8 kHz ogranicza częstotliwość modulacji do poniżej 4 kHz (częstotliwość Nyquista). W oprogramowaniu nie ma jednak jakiegokolwiek filtra antyaliasingowego, co praktycznie ogranicza częstotliwość modulującą do 3 kHz.

Generowanie sygnału wyjściowego

Układ AD9834 jest taktowany z zewnętrznego generatora kwarcowego 80 MHz. Częstotliwość ta jest w tym zastosowaniu niemal optymalna, ponieważ ułatwia filtrowanie w.cz. Generatory kwarcowe 80 MHz są łatwo dostępne w niskich cenach.

Tak jak wspominałem wcześniej, aby wyodrębnić pożądany sygnał w.cz. w paśmie 2 m wymagany jest wysoce selektywny filtr pasmowoprzepustowy o małych stratach. Zastosowano dwie cewki powietrzne o wysokiej dobroci. Ich użycie było warunkiem uzyskania wymaganego rezultatu. Filtr został zaprojektowany tak, aby zapewnić około 40 dB tłumienia najbliższych niepożądanych „aliasów” w pobliżu 100 MHz i 200 MHz i jednocześnie wprowadzać straty nie większe niż 4 dB w paśmie przepustowym. Odpowiednie cewki można na szczęście wykonać szybko i łatwo przy minimalnych kosztach.

Pomimo dużej zmienności poziomów wyjściowych generatora DDS, poziomy sygnału w paśmie 2 m (po filtrze) wynoszą około -40 dBm ± 2 dBm. Przebieg podstawowy generatora, resztkowy sygnał zegara i produkty aliasingu są tłumione o co najmniej 25 dB, a inne składowe pasożytnicze leżą co najmniej 30 dB poniżej poziomu sygnału wyjściowego.

Trzy ustawienia poziomu wyjściowego, przydatne do testowania odbiorników radiowych, zapewnia para przełączanych tłumików. Najwyższy poziom wyjściowy, -45 dBm ± 2 dBm, zapewnia w odbiorniku FM skuteczne ograniczanie amplitudy bez przeciążenia oraz demodulację tonu akustycznego 1 kHz z bardzo dobrym stosunkiem sygnału do szumu.

Załączenie jednego (dowolnego) tłumika skutkuje zmniejszeniem poziomu sygnału do około -75 dBm. Jest to poziom „graniczny” dla typowego odbiornika. Przy tym poziomie ogranicznik w odbiorniku zaczyna korzystnie poprawiać stosunek sygnału do szumu.

Włączenie drugiego tłumika daje poziom sygnału testowego około -105 dBm. Jest to wartość zbliżona do używanej w typowych testach czułości odbiornika „12 dB SINAD” (przy których stosunek sygnału do szumu i zniekształceń wynosi 12 dB; przypis redaktora) oraz w testach układów „squelch” (układy bramki szumów w radiotelefonach; przypis redaktora).

Dokładność bezwzględna tych poziomów zależy częściowo od ustawienia filtra wyjściowego oraz od konstrukcji urządzenia. Prototyp został umieszczony w obudowie wydrukowanej w 3D, zapewniającej jedynie częściowe ekranowanie. Ogranicza to bezwzględną dokładność pomiarów, a zatem dokładność niektórych pomiarów. Jeśli więc zależy Wam na precyzji, musicie użyć ekranującej obudowy metalowej.

Zasilanie

Moim początkowym zamysłem było zasilenie mikrokontrolera ATtiny45 i AD9834 napięciem 5 V ze stabilizatora liniowego 7805.

Karta katalogowa AD9834 podaje, że jest to „układ DDS o małym poborze prądu” (20 mA przy 5 V). Niezbędny jest jednak zewnętrzny generator kwarcowy 80 MHz. Generatory takie pobierają zwykle prąd 30...70 mA (choć niektóre typy zaledwie 10 mA). Dlatego stabilizator małej mocy 78L05 przy napięciu zasilania źródła 9...12 V może się okazać za słaby.

Alternatywnym rozwiązaniem byłoby użycie modułu stabilizatora impulsowego (zdjęcie 1). Można go zamontować na płytce stabilizatora liniowego 7805. Napięcie wyjściowe może wynosić 6...15 V, napięcie wyjściowe to 5 V, a prąd wyjściowy może osiągać 500 mA. Stabilizator impulsowy ma dużą sprawność i nie grzeje się podczas pracy.

Prototyp jest zasilany pojedynczym ogniwem alkalicznym AA 1,5 V. Ze względu na ograniczoną pojemność takiego ogniwa, zasilany nim generator nadaje się tylko do sporadycznego użytku. Możliwe jest jednak wtedy

zastosowanie miniaturowej obudowy PLA wydrukowanej w 3D. Napięcie ogniwa jest podnoszone do 5 V w module przetwornicy podwyższającej (zdjęcie 2).

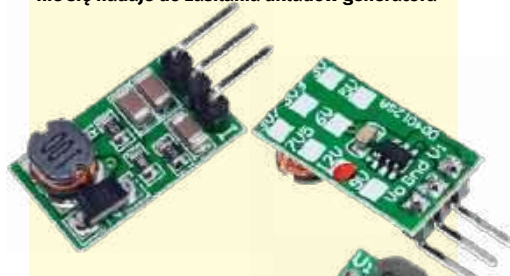
Możliwe jest również zasilanie generatora z pojedynczego ogniwa akumulatora litowo-jonowego (Li-ion) lub litowo-polimernego (LiPo). W stanie pełnego naładowania mają one napięcie 4,2 V i napięcie końcowe 3,5 V. Przetestowałem generator z napięciami zasilania od 3 do 5 V. W całym zakresie napięć poziom wyjściowy pozostawał stały w granicach 0,2 dB!

W przypadku użycia ogniwa Li-ion lub LiPo istnieje możliwość umieszczenia w układzie małej płytki ładowarki USB (zdjęcie 3). Niektóre wersje takich ładowarek zawierają funkcję automatycznego odłączania baterii, by zagwarantować, że bateria nie będzie używana przy napięciu poniżej 3,5 V, co mogłoby ją uszkodzić.

Konstrukcja

Generator jest zbudowany na małej płytce drukowanej o wymiarach 50×70 mm (kod Silicon Chip 06107231). Na płytce znajduje się kombinacja elementów SMD i przewlekanych. Na górnej stronie płytki znajduje się niemal ciągła płaszczyzna masy, a wszystkie elementy SMD są zamontowane

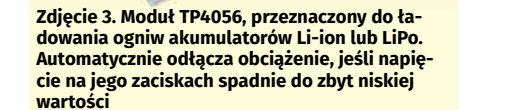
Zdjęcie 1. Mały impulsowy stabilizator obniżający DD4012SA (widok od góry i od spodu), dostarczający napięcie 5 V. Ma dużą sprawność i świetnie się nadaje do zasilania układów generatora



Zdjęcie 2. Niewielki moduł stabilizatora podwyższającego, służący do uzyskania 5 V z pojedynczego ogniwa



Zdjęcie 3. Moduł TP4056, przeznaczony do ładowania ogniw akumulatorów Li-ion lub LiPo. Automatycznie odłącza obciążenie, jeśli napięcie na jego zaciskach spadnie do zbyt niskiej wartości





www.elportal.pl

się od dwukropków. Klikamy myszą z lewej strony pierwszego wiersza (zaraz na prawo od dwukropka), a następnie przeciągamy go w dół do trzeciego wiersza i zwalniamy przycisk myszy. Naciskamy CTRL+C (lub kopiujemy tekst do schowka w inny sposób). Następnie tworzymy nowy plik tekstowy, otwieramy go i wciskamy CTRL+V, wklejając wiersze tekstu. Plik zapisujemy, a następnie zmieniamy jego rozszerzenie z .txt na .eep. Skutek będzie taki sam, jak w Excelu po naciśnięciu zielonego pola.

W każdym przypadku zaleca się zmienić nazwę pliku wynikowego EEP na własną – najlepiej taką, która pomoże nam pamiętać, do czego ten plik służy. Jeśli tego nie zrobimy, to przy następnym użyciu arkusza kalkulacyjnego plik EEP zostanie nadpisany.

Programowanie mikrokontrolera

Plik HEX do generatora, wraz z kodem źródłowym w BASCOM-ie, jest do pobrania na stronie internetowej Silicon Chip. Programując mikrokontroler można jednocześnie załadować do niego plik EEP z czterema zdefiniowanymi częstotliwościami.

Mikrokontroler ATtiny45 programujemy plikami HEX i EEP w odpowiednim programatorze (na przykład USBasp). Następnie programujemy „bity bezpiecznikowe” (konfiguracyjne). Ich wyznaczenie ustawienia pokazano w ramce „Ustawienia bitów bezpiecznikowych ATtiny45”.

Mój artykuł „Shirt Pocket DDS Oscillator” w Silicon Chip z września 2020 r. opisywał mały adapter programujący, który z odpowiednim programatorem szeregowym może być używany do programowania ATtiny w układzie (www.siliconchip.au/Article/14563; str. 47). Płytke drukowaną adaptera można nabyć w sklepie internetowym Silicon Chip (www.siliconchip.au/

Wykaz elementów:

1 dwustronna płytka drukowana, kod 06107231, 50×70 mm
1 obudowa wydrukowana w 3D i nalepka na panel przedni
1 pojemnik ogniwa AA
1 ogniwo alkaliczne AA
1 generator kwarcowy 80 MHz, rozmiar pełny DIP14 lub skrócony (X01) [AliExpress www.siliconchip.au/link/abmb]
3 przyciski do montażu na płytce drukowanej (S1...S3) [Altronics S1126 A, Jaycar SP0609]
4 miniaturowe przełączniki suwakowe DPDT do montażu na płytce drukowanej (S4...S7) [Altronics S2010, Jaycar SS0852]
1 podstawka DIL8 (dla IC1)
1 gniazdo kątowe SMA (CON1)
3 złącza „goldpin”, 2-pinowe, raster 2,54 mm (opcjonalne) (CON2...CON4)
1 3-stykowe złącze kątowe, rozstaw 2,54 mm (jeśli REG1 nie ma takiego złącza)
2 wkręty samogwintujące nr 4, długość 5 mm
2 śruby z łbem walcowym M3, długość 8 mm
2 śruby z łbem stożkowym M3, długość 8 mm
2 nylonowe tuleje dystansowe z gwintem M3, długość 10 mm
40 cm drutu nawojowego miedzianego w emalii, średnica 0,4 mm (26SWG) (dla L1 i L2) [Altronics W0404, Jaycar WW4014]
cienki izolowany przewód potężniejszy

Półprzewodniki:

1 mikrokontroler 8-bitowy ATtiny45-20PU, obudowa DIP-8, zaprogramowany wsadem 0610723A.HEX (IC1)
1 generator DDS AD9834-BRUZ lub -CRUZ, obudowa TSSOP-20 (IC2) [AliExpress www.siliconchip.au/link/abmc]
1 trójkońcówkowa przetwornica podwyższająca 5 V (zasilanie z pojedynczego ogniwa AA) [Silicon Chip SC6780, AliExpress www.siliconchip.au/link/abmd] (REG1)
lub
1 trójkońcówkowa przetwornica obniżająca 5 V (zasilanie z napięcia 6,5...40 V) [Silicon Chip SC6781, AliExpress www.siliconchip.au/link/abme] (REG1)
lub
1 stabilizator liniowy 7805, 5 V 1 A, obudowa TO-220 (zasilanie z napięcia 8...16 V) (REG1)

Kondensatory: (wszystkie ceramiczne SMD M2012/0805 50 V, o ile nie zaznaczono inaczej)

1 szt. elektrolityczny 10 µF 50/63 V, radialny [Altronics R5065, Jaycar RE6075]
1 szt. elektrolityczny 1 µF 50/63 V, radialny [Altronics R5018, Jaycar RE6032]
5 szt. 100 nF X7R
1 szt. 22 nF X7R
1 szt. 10 nF X7R
1 szt. 120 pF NP0/COG
2 szt. 47 pF NP0/COG
2 trymery 6...20 pF do montażu na płytce drukowanej (VC1, VC2) [Altronics R4005]

Rezystory: (wszystkie SMD M2012/0805, 1%)

2 szt. 10 kΩ 1 szt. 6,8 kΩ 1 szt. 3,9 kΩ 1 szt. 1,8 kΩ 2 szt. 820 Ω 4 szt. 51 Ω 1 szt. 220 Ω

Shop/8/5642). Można również używać naszego adaptera do programowania PIC/AVR, opublikowanego w maju i czerwcu 2012 r. (www.siliconchip.au/Series/24) ewentualnie samodzielnie zbudować adapter na płytce prototypowej.

Przeprogramowywanie

Jeśli zechcecie zmienić częstotliwości w pamięci mikrokontrolera, nie będziecie mogli po prostu podłączyć go do zwykłego programatora, ponieważ w zaprogramowanym mikrokontrolerze pin RESET jest wyłączony. ATtiny45 należy najpierw skasować

w specjalnym programatorze używającym podwyższonego napięcia programującego.

Zaprojektowałem proste urządzenie do kasowania układów i przywracania stanu bitów bezpiecznikowych – CEFR („Chip Eraser and Fuse Restorer”). O tym, jak je zbudować, można przeczytać na mojej stronie internetowej pod adresem www.zl2pd.com/CEFR_Fuse_Reset_Tool.html. CEFR nie wymaga żadnych specjalnych elementów. Może być zasilany z gniazda USB lub z zewnętrznego zasilacza USB 5 V. Innym znanym amatorskim narzędziem do kasowania bitów bezpiecznikowych jest Fuse Doctor (<https://github.com/SukkoPera/avr-fusebit-doctor>).

Montaż końcowy

W zależności od wybranych wariantów obudowy i zasilania należy dodać

Dostępność elementów

Do tego projektu nie powstał zestaw do samodzielnego montażu. Silicon Chip może jednak dostarczyć płytkę drukowaną, zaprogramowany mikrokontroler i moduł przetwornicy obniżającej lub podwyższającej. Pozostałe elementy można znaleźć u wielu dostawców, na przykład tych wymienionych w liście elementów.



Zdjęcie 4. Prosty „Chip Eraser and Fuse Restorer” (CEFR), który kasuje bity bezpiecznikowe w ATtiny25/45/85 do stanu fabrycznego, dzięki czemu można w pamięci EEPROM przeprogramować częstotliwości użytkownika



Zdjęcie 5. Płytkę drukowaną mieści się w obudowie wydrukowanej w 3D i jest zasilana z pojedynczego ogniwa AA 1,5 V. Moduł przetwornicy podwyższającej znajduje się po lewej stronie przełącznika zasilania. Filtr wyjściowy jest umieszczony powyżej i na lewo od stabilizatora. AD9834 i inne elementy SMD znajdują się na spodzie płytki drukowanej, ale ich położenia są zaznaczone na górnej stronie płytki

wyłącznik zasilania i okablować go. W prototypie ogniwo AA znajduje się w dolnej części obudowy zaprojektowanej specjalnie dla generatora i wydrukowanej w 3D. Pliki STL tej obudowy są do pobrania tam, gdzie oprogramowanie – na stronie www.siliconchip.com.au/Shop/6/266.

Pojemnik baterii montujemy, wsuwając jego zaciski w przeznaczone dla nich szczeliny w obudowie. Wyłącznik zasilania również wsuwamy w przeznaczone na niego miejsce.

Następnie montujemy płytkę drukowaną. Dwa samogwintujące wkręty o długości 5 mm i średnicy 3 mm wchodzą od spodu w dwie nylonowe tuleje dystansowe o długości po 6 mm, utrzymując płytkę w miejscu. Kolejne cztery śrubki mocują od góry panel przedni.

Testowanie i działanie

Wkładamy zaprogramowany mikrokontroler ATtiny45 do podstawki, upewniając się, że jego pin 1 znajduje się przy oznaczonym końcu podstawki. W razie wątpliwości sprawdzamy położenie mikrokontrolera według rysunku 7.

Przełączamy oba przełączniki tłumika w skrajne lewe położenia (tłumienie minimalne).

Umieszczamy przenośny odbiornik FM pasma 2 m w odległości około 50 cm

od zmontowanej płytki drukowanej i nastawiamy go na odbiór częstotliwości 146,000 MHz. W odbiorniku wyłączamy wysłuchanie aby słyszeć szum kanału. Ustawiamy odpowiednią siłę głosu.

Załączamy zasilanie generatora. Szum w odbiorniku FM powinien natychmiast ucichnąć, a miernik siły sygnału (jeśli odbiornik go zawiera) powinien wskazywać poziom bardzo wysoki (zazwyczaj S9 lub lepszy). Jeśli tak się nie dzieje, wyłączamy zasilanie generatora i dokładnie sprawdzamy cały jego układ.

Jeśli generator ten pierwszy test przeszedł pomyślnie, naciskamy na chwilę i puszczaamy przycisk modulacji. W głośniku odbiornika powinniśmy usłyszeć czysty ton 1 kHz.

Kontrolując dźwięk i obserwując miernik poziomu sygnału odbiornika, regulujemy dwoma kondensatorami dostrojczy generatora, aż uzyskamy maksymalny poziom sygnału. Może być konieczne odsunięcie odbiornika nawet o kilka metrów, tak aby w odbieranym sygnale zaczął się już pojawiać niewielki szum.

Niewielką poprawę dostrojenia można czasem uzyskać poprzez bardzo nieznaczne ściśnięcie lub rozciągnięcie jednej lub obu cewek, ale rzadko jest to konieczne.

Generator uruchamiamy się z częstotliwością ustawioną na 146,0 MHz, bez modulacji. Każde naciśnięcie przycisku FREQ („częstotliwość”) zwiększa częstotliwość o 500 kHz. Jeśli bieżąca częstotliwość wynosi

148,0 MHz, następne naciśnięcie przycisku zmieni ją na 144,0 MHz, a kolejne naciśnięcia znów będą ją zwiększać o 500 kHz.

Kontynuując test, naciskamy przycisk FREQ, aby zmienić częstotliwość generatora na 146,5 MHz. Dostrajamy odbiornik do tej częstotliwości. Powinien być słyszalny modulowany sygnał.

Naciskając krótko przycisk MOD („modulacja”) sprawdzamy, czy modulacja się załącza i wyłącza. Przed przejściem do następnego punktu testu musimy się upewnić, że modulacja jest załączona.

Po naciśnięciu przycisku MEM („pamięć”) generator wytworzy jedną z czterech zaprogramowanych częstotliwości, zaczynając od pamięci nr 1. Naciśnięcie przycisku FREQ spowoduje powrót generatora do częstotliwości ostatnio wybranej tym przyciskiem w krokach 500 kHz.

Stroimy odbiornik na częstotliwość zaprogramowaną w pamięci 1 i wybieramy tę pamięć, naciskając krótko przycisk MEM. W tym kanale powinien się teraz pojawić modulowany sygnał. Nastawiając częstotliwość przyciskiem FREQ lub wybierając ją z pamięci, testujemy inne częstotliwości generatora.

Istotne wskazówki

Przy pewnych ustawieniach częstotliwości, na wyjściu generatora będą się pojawiać dodatkowe składowe, leżące w paśmie 2 m lub poza nim. Po tak prostym generatorze DDS

Skąd tu modulacja fazy?

Być może zauważyliście, że wciąż piszę o modulacji fazy (PM), natomiast w opisie projektu jest mowa o trybie FM (modulacji częstotliwości). O co chodzi?

Modulacja fazy i modulacja częstotliwości to jakby dwie strony tej samej monety – „modulacji kąta”.

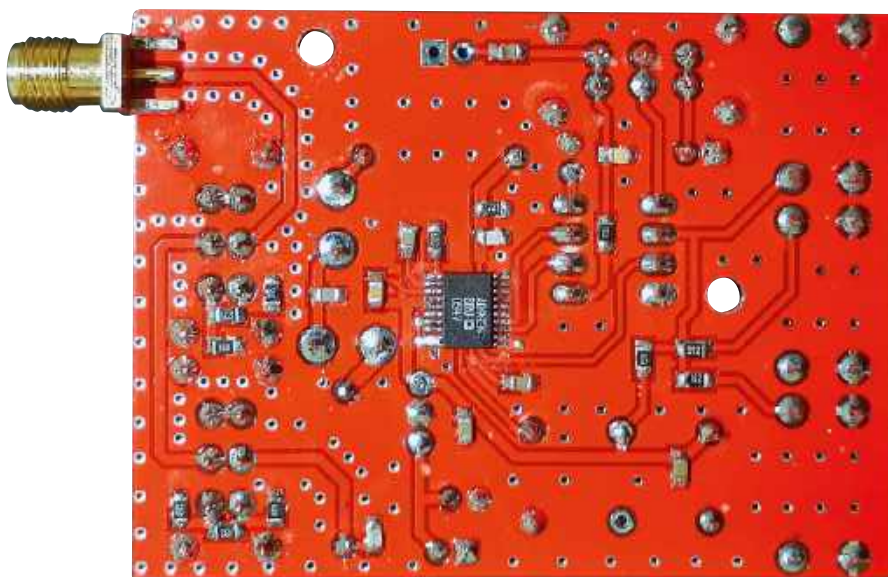
Jeśli modulator jest rzeczywiście modulatorem fazy, to w celu zrealizowania modulacji FM modulujący sygnał akustyczny powinien przed modulacją zostać scałkowany. Analogicznie – jeśli mamy modulator częstotliwości, a chcemy uzyskać modulację fazy, to sygnał modulujący musimy zróżniczkować. **Przypis redaktora: całkowanie i różniczkowanie sygnału można realizować w znanych układach ze wzmacniaczami operacyjnymi, albo, w przybliżeniu, w filtrach RC – odpowiednio dolnoprzepustowym i górnoprzepustowym.**

Wszyscy radioamatorzy mówią o „FM” niezależnie od tego, jak naprawdę rodzaj modulacji jest stosowany w ich sprzęcie. W nadajnikach i innych układach z analogową modulacją kąta zawsze łatwiej było mierzyć dewiację częstotliwości. Nikt nigdy nie mierzył przesunięcia fazowego. Tak więc, gdy w grę wchodzi analogowa modulacja kąta, wszyscy mówią o FM i dewiacji częstotliwości.

W opisywanym tu generatorze testowym modulację FM i PM byłoby nie do odróżnienia, ponieważ używany jest pojedynczy ton. **Przypis redaktora: ton czyli sygnał sinusoidalny pozostaje sinusoidalnym zarówno po scałkowaniu jak i różniczkowaniu, tyle że ulega przesunięciu w fazie o stały kąt, odpowiednio +90° lub -90°.**

Co ciekawe (ale logiczne) – mówimy o określonych chwilowych przesunięciach fazy (i mierzymy je) w przypadku modulacji kąta w celu transmisji danych cyfrowych. Stąd biorą się systemy PSK, QPSK i 8PSK, gdzie znajduje zastosowanie modulacja o stosunkowo dużym kroku fazowym.

Być może pamiętacie, że wspominałem o możliwości modulacji amplitudy (AM). Niewiele osób używa AM w paśmie 2 m, ale – co ciekawe – połączenie AM i PM umożliwia przy niewielkim dodatkowym wysiłku zrealizowanie modulacji kwadraturowej 16-QAM. Ale to już pomysł na inny dzień...



Zdjęcie 6. Spód płytki drukowanej z zamontowanymi wszystkimi układami SMD. Piny 9 i 10 układu AD9834 są zwarte cyną, ale oba są dołączone do masy, więc cyny nie usuwałem

można się było tego spodziewać. Analog Devices ostrzega o tym w karcie katalogowej.

Te dodatkowe sygnały mogą pochodzić od zegara DDS, jego harmonicznych lub stanowią produkty mieszania się jednego lub więcej „aliasów”. Zazwyczaj leżą one co najmniej 25 dB („aliasy”) lub 30 dB (inne) poniżej znamionowego poziomu sygnału wyjściowego. Są to parametry takie jak w niektórych fabrycznych generatorach w.cz. starszej konstrukcji.

Należy brać pod uwagę brak filtra antyaliasingowego na wejściu modulującym i unikać pokusy podawania na to wejście wzmacnionego dźwięku z mikrofonu. Nie należy również dodawać na wyjściu generatora żadnego wzmacniacza w.cz. Generator nie jest przeznaczony do użycia jako nadajnik FM na pasmo 2 m. Nadaje się on do przeprowadzania podstawowych testów, ale nie spełnia żadnych norm wymaganych dla nadajników.

Modulacja zewnętrzna

Do modulowania sygnału generatora można użyć zewnętrznego źródła dźwięku na przykład generatora o częstotliwości akustycznej. Źródło można dołączyć do wejścia CON2 na płytce. Musimy się upewnić, że poziom międzyszczytowy tego sygnału leży w zakresie od 0 do wartości napięcia zasilania (czyli maksymalnie 5 Vpp przy zasilaniu 5 V). Silniejsze sygnały mogą spowodować uszkodzenie mikrokontrolera. Sygnał modulujący powinien mieć składową stałą około 2,5 V. Słyszalność tonu modulującego kontrolujemy na bieżąco w odbiorniku.

Jeśli chcecie wyprowadzić gniazdo zewnętrznego sygnału modulacji, proponujemy

podłączenie tego gniazda do złącza CON2 poprzez kondensator 100 nF i dołączenie zacisku sygnałowego tego złącza do dzielnika napięcia 5 V, utworzonego z dwóch rezystorów 100 kΩ.

CTCSS

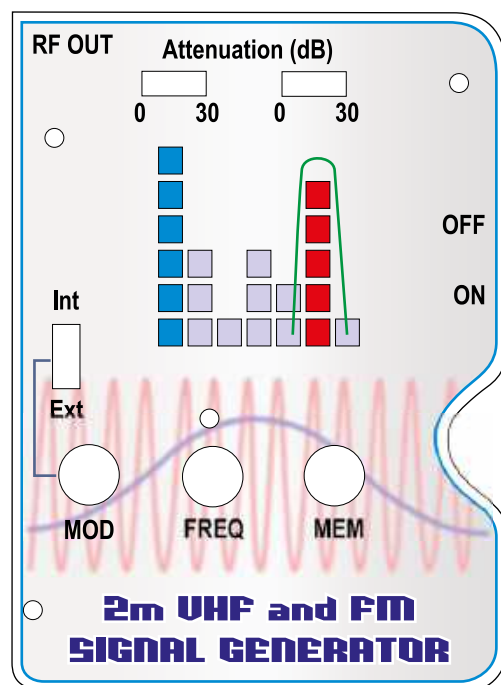
Do wejścia CON4 („CTCSS IN”) na płytce można podłączyć koder CTCSS i wykorzystać go na przykład do przetestowania dekodera CTCSS w odbiorniku FM.

Aby przeprowadzić taki test, ustawiamy ręcznie odpowiednią częstotliwość fali nośnej generatora, a dekodek CTCSS w odbiorniku – na używaną częstotliwość CTCSS, na przykład 123 Hz. Sprawdzamy, czy maksymalny poziom wyjściowy kodera CTCSS nie przekracza około 0,5 Vpp. W generatorze wybieramy wewnętrzne źródło modulacji FM (przełącznik S4 w pozycji „INT”; przypis redaktora) i, aby ją włączyć, naciskamy krótko przycisk S1. Uruchamiamy koder CTCSS. Dekoder CTCSS w odbiorniku powinien wykryć ton CTCSS, wyłączyć wyciszenie i powinien być słyszalny ton 1 kHz z generatora.

Wyłączenie kodera lub zmniejszenie jego poziomu poniżej około 0,1 Vpp powinno spowodować, że dekodek CTCSS w odbiorniku wyciszy dźwięk. Niektóre dekodery mogą być jednak bardzo czułe i wciąż wykrywać prawidłowy ton CTCSS przy poziomie wynoszącym zaledwie 5 mVpp!

Uwagi końcowe

Trudności, na jakie napotkałem w znajdowaniu zarówno podstawowych informacji dotyczących modulacji fazy w generatorach DDS jak i przykładów jej zastosowań, zaskoczyła



Rysunek 8. Naklejka na panel przedni. Można ją wydrukować, załaminować i przykleić do górnej części obudowy

mnie, biorąc pod uwagę, że karty katalogowe mocno tę funkcję promują. Osiągnięcie dokładnych poziomów modulacji fazy (i częstotliwości) w AD9834 okazało się jednak stosunkowo proste.

Mam nadzieję, że projekt jako całość oraz jego szczegóły realizacyjne okażą się interesujące. Miniaturowy generator sygnału FM jest bardzo wygodny w użyciu i stanowi częsty temat rozmów w gronie znajomych fachowców. ■

Andrew Woodfield ZL2PD



Materiały dodatkowe dostępne są na stronie: <https://www.siliconchip.com.au/Shop/8/6820>

Materiały dodatkowe są również dostępne na stronie elportal.pl/do-pobrania

Artykuł reprodukowano na podstawie umowy z magazynem „Silicon Chip”, 2022. www.siliconchip.com.au

Bezprzewodowy Internet o globalnym zasięgu

Spółka Starlink należąca do firmy SpaceX jest dostawcą Internetu o globalnym zasięgu. Zdalna łączność dla urządzeń Internetu Rzeczy (IoT) znajdujących się w dowolnym miejscu na Ziemi może być także realizowana za pośrednictwem innej spółki zależnej, takiej jak Swarm, natomiast usługi oferowane przez spółkę Starshield są przeznaczone dla organizacji rządowych i wojskowych.



Większość mieszkańców krajów rozwiniętych, a nawet rozwijających się, może obecnie korzystać z Internetu za pośrednictwem smartfonów lub innych urządzeń mobilnych, znajdujących się w pobliżu miast lub terenów zamieszkałych, a także wzdłuż głównych tras komunikacyjnych. Łączność z Internetem poza takimi obszarami jest zazwyczaj droga i powolna.

System Starlink jest własnością firmy SpaceX i umożliwia zwykłym użytkownikom przystępny cenowo dostęp do Internetu satelitarnego w dowolnym miejscu na świecie, bez względu na to, czy znajdują się na pełnym morzu, w samolocie, na Antarktydzie czy podczas innej wyprawy na odludzie.

Małe opóźnienia

Oprócz dostawy przystępnych cenowo usług o globalnym zasięgu, Starlink zapewnia małe opóźnienia podczas transmisji danych internetowych w obu kierunkach, czyli do i od użytkownika. Fale radiowe poruszają się z prędkością światła (około 3×10^8 m/s), czyli na pokonanie dystansu między nadajnikiem a odbiornikiem potrzebują czasu. Istnieją również dodatkowe opóźnienia wynikające z przetwarzania sygnału radiowego i dystrybucji danych w Internecie.

Realny czas biegu fal radiowych do satelity geostacjonarnego, orbitującego na wysokości 35 786 km, i z powrotem na Ziemię wynosi około 600 ms lub więcej. Wynikające z tego opóźnienie jest zbyt duże dla dwukierunkowego przekazu dźwięków lub obrazów w czasie rzeczywistym, a także obsługi gier lub innych aplikacji interaktywnych.

System Starlink zapewnia małe opóźnienia dzięki umieszczeniu swoich satelitów na niskich orbitach, na wysokości około 550 km. W ten sposób czas transmisji nie przekracza 20 ms i jest porównywalny z wartościami uzyskiwanymi w sieciach przewodowych. Jednakże, ze względu na tak niskie orbity, dla uzyskania globalnego zasięgu niezbędne jest zastosowanie znacznej liczby satelitów.

Źródło obrazu ze strony SpaceX (CC BY-NC 2.0):
www.flickr.com/photos/spacex/49422067976/in/photostream/

Innym deklarowanym celem budowy systemu Starlink jest zapewnienie łączności internetowej w krajach rozwijających się, w których przewodowa lub bezprzewodowa infrastruktura telekomunikacyjna jest słabo rozwinięta. Według ONZ około 57% światowej populacji nie ma dostępu do Internetu.

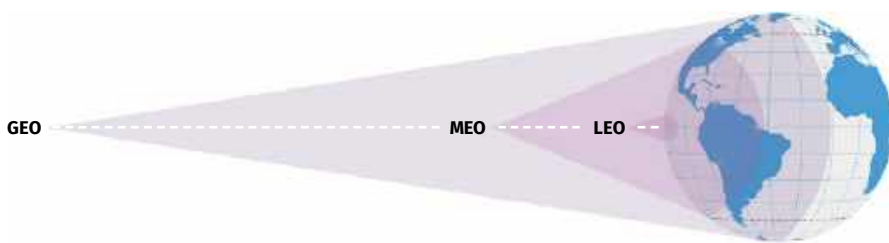
Firma SpaceX

Firma SpaceX, czyli Space Exploration Technologies Corp, jest w dużej części własnością koncernu Elon Musk Trust (47,4% udziałów, 78,3% głosów). SpaceX buduje satelity Starlink, Swarm i Starshield oraz systemy do ich wynoszenia na orbitę, takie jak rakiety Falcon 9. W grudniu 2022 r. Starlink miał milion klientów, w tym znaczną część w Australii i Nowej Zelandii.

Konstelacje satelitów i ich orbity

Konstelacja to grupa satelitów współpracujących ze sobą jako zintegrowany system. Dobrze znanym przykładem jest konstelacja satelitów GPS. Systemy Starlink, Swarm i Starshield również wykorzystują konstelacje satelitów.

Ze względu na niskie opóźnienia wymagane w systemie Starlink, satelity muszą znajdować się na niskich orbitach okołoziemskich. Z tego powodu widoczność



Rysunek 1. Trzy popularne typy orbit i porównanie naziemnych obszarów pokrycia: orbita geosynchroniczna (GEO), średnia orbita okołoziemska (MEO) i niska orbita okołoziemska (LEO)



Rysunek 2. Porównanie cech kilku wysokości orbitalnych. RTT to czas podróży sygnału radiowego w obie strony. Najlepiej unikać pasów radiacyjnych Van Allena. Promień jest liczony od środka Ziemi, natomiast wysokość jest liczona od powierzchni Ziemi. Źródło: <https://www.wiki/6H8X>

Tabela 1. Charakterystyka różnych orbit satelitarnych

	Orbita geosynchroniczna (GEO)	Średnia orbita okołoziemska (MEO)	Niska orbita okołoziemska (LEO)
Wysokość	35 790 km (typowo)	2 000...35 786 km (20 500 km typowo)	160...2 000 km (500 km typowo)
Przykłady	GOES, Inmarsat, Intelsat	GPS (20 180 km)	Starlink (550 km)
Opóźnienie przy typowej wysokości (RTT)	600 ms	400 ms	20 ms
Udział widocznej powierzchni Ziemi	42,4%	38,1%	4,0%
Minimalna liczba satelitów dla pokrycia całej Ziemi	Trzy; cztery dla częściowego nakładania się zasięgów	10...15; więcej dla redundancji	Co najmniej 32, ale w praktyce setki
Wymagana prędkość śledzenia anteny	Bez śledzenia	Wolne; każdy satelita widoczny przez 1...3 godziny	Szybkie; każdy satelita widoczny przez 5...15 minut
Zalety	Wystarczy niewiele satelitów, brak śledzenia, brak przełączania (handover), stałe połączenie, proste zarządzanie, brak skomplikowanych orbit	Mniej satelitów niż w LEO, mniejsze opóźnienia niż w GEO; mniejsze systemy antenowe; lepsza siła sygnału powyżej 72° szerokości geograficznej	Małe opóźnienia, niewielkie tłumienie sygnału, stacje naziemne o małej mocy, potencjalnie niższy koszt dzięki produkcji masowej
Wady	Słaby sygnał, słabe pokrycie powyżej 72° szerokości geograficznej, duże opóźnienia	Wymagane śledzenie anteny, konieczne przełączanie satelitów, więcej satelitów niż w GEO, większe narażenie na promieniowanie pasów Van Allena	Mały obszar obsługi, wymagane śledzenie anteny, częste przełączanie satelitów, duże przesunięcia Dopplera, krótki czas życia orbity z powodu oporu atmosfery
Typowy czas życia satelity	>15 lat	10...15 lat	3...7 lat
Złożoność sieci	Niska	Średnia	Wysoka

pojedynczego satelity dla obiektu umieszczonego na Ziemi jest ograniczona. Dlatego do pełnego pokrycia całej kuli ziemskiej wymagane jest użycie znacznej liczby satelitów.

Na **rysunku 1** przedstawione jest porównanie obszaru pokrycia na powierzchni Ziemi dla trzech popularnych wysokości orbitalnych: orbity geosynchronicznej (GEO), średniej orbity okołoziemskiej (MEO) i niskiej orbity okołoziemskiej (LEO). Satelity Starlink są umieszczane na orbicie LEO. Na **rysunku 2** i w **tabeli 1** przedstawione są dane dotyczące orbit GEO, MEO i LEO.

Obszar na powierzchni Ziemi widoczny dla satelity umieszczonego na określonej wysokości można wyznaczyć z poniższego wzoru:

$$f = d \div 2 \times (R + d)$$

gdzie d oznacza wysokość orbity nad powierzchnią Ziemi, R oznacza promień Ziemi (nominalnie 6378 km na równiku), zaś f oznacza ułamek powierzchni Ziemi widoczny z satelity. Ten wzór został użyty do wyznaczenia wartości umieszczonych w tabeli 1.

Ze względów praktycznych żaden z satelitów nie będzie widoczny aż po horyzont, z powodu obecności gór, drzew itp. Ponadto, sygnał radiowy będzie degradowany przez długą drogę przez atmosferę. Dlatego w praktyce definiuje się pewien kąt elewacji, poniżej którego nie podejmuje się prób komunikacji między satelitą a stacją naziemną (**rysunek 3**).

Czerwony dysk reprezentuje absolutny horyzont, zaś żółty oznacza minimalny zasięg przy zadanych kącie elewacji. W związku z tym zasięg jaki może osiągnąć satelita jest mniejszy niż obliczony za pomocą powyższego wzoru.

Wysokość orbitalna satelitów Starlink

Jak w każdym dużym programie satelitar-
nym, przewidziano użycie satelitów w róż-
nych wersjach. Obecnie w przypadku Starlink
produkowane są satelity V1, V1.5 (**rysunek 4
i 5**), V2 i V2 mini. Do tej pory na orbitę trafiły
tylko satelity V1 i V1.5, przy czym satelity V1.5
są nadal wystrzeliwane (**obecnie na orbitach
pracują satelity w wersjach V2 i V3 – stan
na styczeń 2026 r. – przypis red.**).

Satelity Starlink V1 są rozmieszczone w różnych powłokach, na orbitach o wysokości

od 540 km do 570 km, pokazanych w **tabeli 2**. Powłoka orbitalna to grupa satelitów dzielących tę samą orbitę kołową na określonej wysokości.

Harmonogram rozmieszczania satelitów

Firma SpaceX stale wyszłeuje nowe satelity Starlink. W 2018 r. wyszłeuono dwa satelity testowe „Tintin”. W 2019 r. wyszłeuono kolejną serię 60 satelitów V0.9 należących do projektu produkcyjnego.

Firma SpaceX wyrzelała satelity V1.0 od listopada 2019 do maja 2021 roku. Zwykle jednorazowo wyrzelało 60 satelitów. Po 29 startach, sumaryczna liczba wyniosła 1675 satelitów. W chwili obecnej 183 z nich już nie działa.

Wyrzeliwanie satelitów Starlink V1.5 rozpoczęło się w czerwcu 2021 roku i trwało do stycznia 2023 roku lub dłużej. Do tej pory odbyło się 40 startów rakiet z satelitami V1.5, z których każda wynosiła do 54 satelitów. W sumie zostało wyrzelandonych 1881 satelitów, z czego 52 już nie działa.

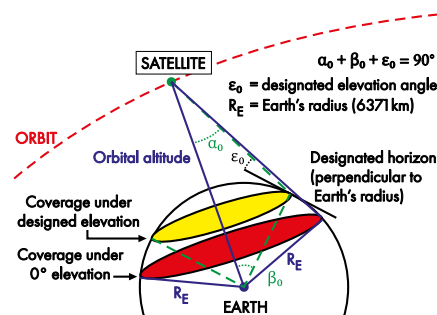
Pełna, aktualna lista danych dotyczących satelitów Starlink znajduje się na stronie <https://planet4589.org/space/con/star/stats.html>.

13 stycznia 2022 roku odbyły się cztery starty rakiet z satelitami Starshield V1.5. Kolejne cztery odbyły się 19 czerwca 2022 roku, z przeznaczeniem dla nieznanych agencji rządowych USA.

Obszar pokrycia

Początkowo, dzięki pierwszej powłoce orbitalnej, na orbicie nachylonej pod kątem $53,0^\circ$, system Starlink zapewniał zasięg na obszarach poniżej około 55° szerokości geograficznej, co obejmuje zdecydowaną większość światowej populacji. Późniejsze uruchomienia na innych inklinacjach orbitalnych pokrywały wyższe szerokości geograficzne.

Powłoka 53,2° rozszerzyła liczbę klientów objętych zasięgiem na średnich i niskich szerokościach geograficznych. Powłoka 70,0° rozszerzyła zasięg na Alaskę i północną Europę (i na równoważne szerokości geograficzne na półkuli południowej).



Rysunek 3. Widoczność satelity przy zerowej wysokości (czerwony) i projektowanej wysokości (żółty), pokazująca różnicę między teoretycznym a rzeczywistym zasięgiem. źródło: www.frontiersin.org/articles/10.3389/frcmn.2021.643095/full

Wcześniej szerokie rozmieszczenia satelitów Starlink odbywały się na orbitach równikowych, więc nie pokrywały regionów polarnych (**rysunek 6**).

Cztery starty wynoszące na orbitę po 46 satelitów dla powłoki 97,6° miały miejsce w lipcu i sierpniu 2022 roku, dodając pokrycie dla regionów polarnych. Obejmuje to Antarktydę oraz obszary północnej Alaski, północnej Kanady, Finlandii, Norwegii i Szwecji, które wcześniej nie były objęte zasięgiem.

Wysokopoziomowa architektura Starlink

System Starlink składa się z trzech głównych elementów: satelitów, stacji naziemnych i terminali użytkowników. Stacje naziemne zapewniają połączenie z naziemną siecią internetową i mogą stanowić środek komunikacji między satelitami.

Liczba potrzebnych stacji naziemnych została zminimalizowana dzięki wprowadzeniu najnowszych satelitów z serii V1.5+, które mogą komunikować się między sobą z użyciem łączy laserowych. Gdy użytkownik łączy się z satelitą za pośrednictwem swojego terminala, satelita Starlink przekazuje jego sygnał do stacji naziemnej lub do innego satelity Starlink, a dopiero później do stacji naziemnej.

Przekazywanie sygnału między-satelitami jest niezbędne do obsługi użytkowników znajdujących się na dużych szerokościach geograficznych, gdzie satelity mają dostęp do niewielu stacji naziemnych lub nie mają go wcale.

Tabela 2. Powłoki orbitalne i liczba satelitów Starlink V1 i V1.5 (łącznie 4408)

Powłoka	Nachylenie orbity	Wysokość orbity	Liczba płaszczyzn orbitalnych	Docelowa liczba satelitów/płaszczyznę	Łączna liczba satelitów	Łączność laserowa
Powłoka 1	53,0°	550 km	72	22	1584	
Powłoka 2	70,0°	570 km	36	20	720	
Powłoka 3	97,6°	560 km	6 (polarne)	58	348	Wszystkie
Powłoka 4	53,2°	540 km	72	22	1584	
Powłoka 5	97,6°	560 km	4 (polarne)	43	172	Wszystkie

Komunikacja laserowa między satelitami

Satelity V1.5 mogą komunikować się między sobą za pośrednictwem łącz laserowych. Zmniejsza to opóźnienia, ponieważ sygnał przesyłany za pośrednictwem lasera będzie podróżował o około 30...40% szybciej niż sygnał przesyłany między urządzeniami naziemnymi, połączonymi za pomocą kabla koncentrycznego lub światłowodu.

Ponadto, ze względu na krótszy czas transmisji danych między satelitami niż między stacjami naziemnymi, połączonymi za pomocą kabli naziemnych lub podmorskich, ogólne opóźnienia są zmniejszane nawet o 50%.

Laserowe połączenia między satelitami są szczególnie ważne dla satelitów Starlink znajdujących się na orbicie polarnej, ponieważ nie będą one miały dostępu do stacji naziemnych.

System Starlink

System Starlink może być używany na całej powierzchni Ziemi, jednak zgodnie z przepisami Międzynarodowego Związku Telekomunikacyjnego (ITU) i traktatami międzynarodowymi, w każdym z krajów muszą być przyznane osobne prawa do korzystania z komunikacji satelitarnej, takiej jak Starlink.

Nowa Zelandia i Australia w kwietniu 2021 roku zatwierdziły możliwość użytkowania systemu Starlink na swoim terenie. Stały się piątym i szóstym krajem po USA, Kanadzie, Wielkiej Brytanii i Niemczech, które jako pierwsze wyraziły taką zgodę. Na **rysunku 7** przedstawiono dostępność systemu Starlink w poszczególnych krajach.

Urządzenia Starlink objęte planem mieszkaniowym są przystosowane do pracy tylko pod adresem użytkownika lub w pobliżu jego domu, a także na innych obszarach objętych planem RV.

Liczba satelitów na orbicie

Aby docenić ogrom projektu Starlink, należy wziąć pod uwagę liczbę satelitów już znajdujących się na orbicie. Według indeksu Biura Narodów Zjednoczonych do spraw Przestrzeni Kosmicznej (UNOOSA) dostępnego na stronie www.unoosa.org/oosa/oso-index/search-ng.jspx, na dzień 3 stycznia 2023 roku, od czasów Sputnika 1 w kosmos zostało wystrzelonych 14281 obiektów. Spośród nich 8734 sklasyfikowano jako nadal przebywające na orbicie, choć niekoniecznie działające.

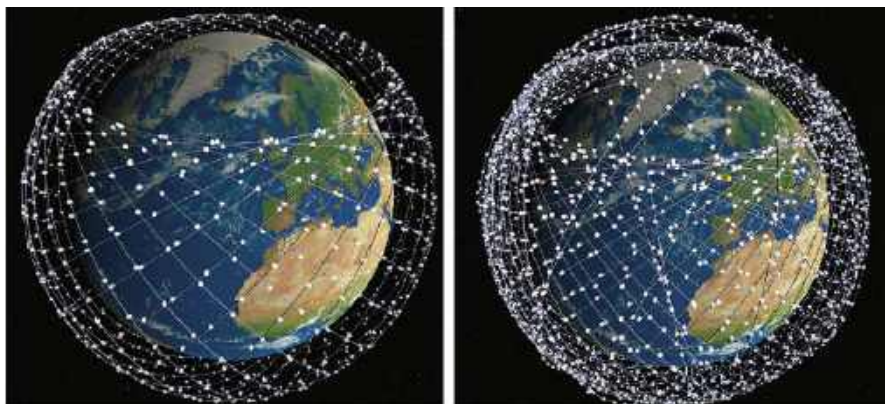
Spośród 8734 obiektów sklasyfikowanych jako umieszczone na orbicie, 3568 było należących do systemu Starlink, a 5166 nie.



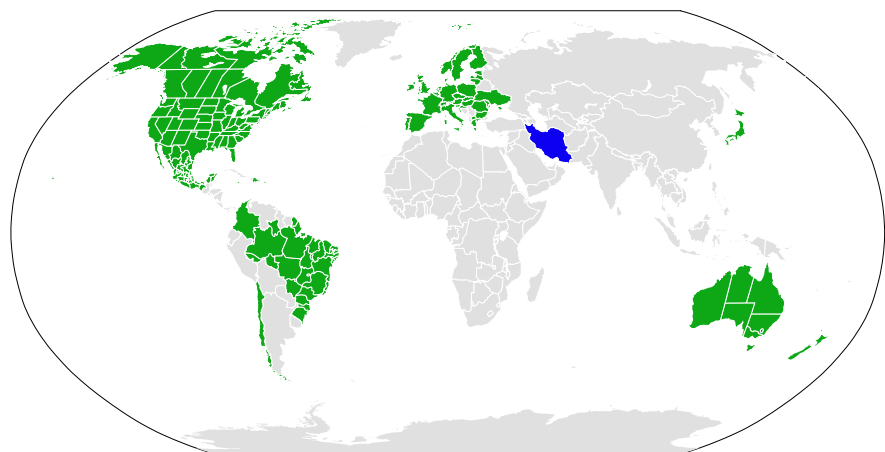
Rysunek 4. Artystyczna koncepcja satelity StarlinkV1.1. Źródło: <https://w.wiki/6H8Y>



Rysunek 5. Renderowanie satelitów StarlinkV1.5 (po lewej) i V1 (po prawej). Satelity V2.0 mają pięciokrotnie większą powierzchnię anten skierowanych na Ziemię i są znacznie bardziej wydajne. Źródło: www.teslarati.com/spacex-elon-musk-next-gen-starlink-satellite-details/



Rysunek 6. Niepełny zasięg globalny zapewniony przez wcześniej wystrzelone satelity Starlink na orbitach równikowych (po lewej) w porównaniu do pełnego zasięgu globalnego po późniejszych wystrzeleniach na orbitach polarnych (po prawej)



Rysunek 7. Dostępność usług Starlink. Zielony oznacza obszar zatwierdzony i aktywowany, niebieski oznacza aktywowany, a szary jest nieznan. Źródło: <https://w.wiki/6H8Z>

Oznacza to, że prawie 41% orbitujących obiektów jest powiązanych z systemem Starlink. Liczba ta wzrośnie diametralnie wraz z rozwojem całej konstelacji. Tak więc za kilka lat znaczna większość wszystkich sztucznych satelitów Ziemi może być częścią systemu Starlink!

Według danych na dzień 20 stycznia 2023 roku, pochodzących ze strony internetowej <https://planet4589.org/space/con/star/stats.html>, która prowadzi rejestr satelitów Starlink (stan na dzień 11 stycznia 2026 roku: wystrzelonych urządzeń – 10897, pracujących – 9440 – przypis red.).

Cechy wyróżniające satelity Starlink

Do innych cech satelitów Starlink, o których jeszcze nie wspomniano, należą:

- płaski kształt dla łatwiejszego i gęstszego upakowania satelitów do rakiet Falcon 9
- nawigacja na podstawie położenia gwiazd
- satelity są wyposażone w cztery anteny sterowane fazowo i dwie paraboliczne (www.starlink.com/technology).

Wygląd abonenckich anten Starlink przeznaczonych dla indywidualnych użytkowników pokazano na rysunku 8.

System Starlink w lotnictwie cywilnym

Co prawda połączenie z Internetem było już wcześniej dostępne w niektórych samolotach, jednak miało ono niską przepustowość i było kosztowne. System Starlink umożliwia użytkownikom znajdującym się na pokładzie samolotu korzystanie z usług wymagających połączeń o dużej przepustowości i niskich opóźnieniach. Pasażerowie mogą wykonywać pracę zdalną podczas lotu, oglądać filmy o wysokiej rozdzielczości, grać gry online itp. Urządzenia abonenckie uzyskują dostęp

Częstotliwości używane przez satelity Starlink

Według www.elonx.net/starlink-compendium/, następujące częstotliwości są używane przez Starlink:

- **Satelita do terminali użytkownika:** 10,7...12,7 GHz, 37,5...42,5 GHz
- **Satelita do bramki:** 17,8...18,6 GHz, 18,8...19,3 GHz, 37,5...42,5 GHz
- **Terminale do satelitów:** 14,0...14,5 GHz, 47,2...50,2 GHz, 50,4...51,4 GHz
- **Bramy do satelitów:** 27,5...29,1 GHz, 29,5...30,0 GHz, 47,2...50,2 GHz, 50,4...51,4 GHz
- **Śledzenie, telemetria i kontrola (łączy w dół):** 12,15...12,25 GHz, 18,55...18,60 GHz, 37,5...37,75 GHz
- **Śledzenie, telemetria i kontrola (uplink):** 13,85...14,00 GHz, 47,2...47,45 GHz

do Internetu za pośrednictwem standardowego łącza Wi-Fi.

Rysunek 9 przedstawia antenę z certyfikatem lotniczym przeznaczoną do pracy w systemie Starlink. Związana z tym usługa została uruchomiona w 2023 roku. Szybkość łącza wynosi 350 Mb/s a opóźnienia w transmisji nie przekraczają 20 ms.

Dla osób zainteresowanych kosztami, w chwili pisania tego tekstu koszt sprzętu wynosi 150 000 USD, a miesięczne opłaty za usługi z nieograniczoną ilością danych wynoszą 12 500...25 000 USD.

Wstępna certyfikacja jest uzyskana dla następujących typów samolotów biznesowych i regionalnych: ERJ-135, ERJ-145, G650, G550, Falcon 2000, G450, Challenger 300, Challenger 350, Global Express, Global 5000, Global 6000 i Global 7500, przy czym opracowywane są kolejne aplikacje dla większych, komercyjnych odrzutowców.

Jak działają anteny Starlink

W przeciwieństwie do anten skierowanych na satelitę geostacjonarnego, które wymagają dobrej widoczności tylko w jednym kierunku, anteny Starlink muszą odbierać sygnały od horyzontu do horyzontu. Wynika to z faktu, że satelity LEO mogą znajdować się w dowolnym miejscu na niebie. Aplikacja uruchamiana na telefonie komórkowym pozwala na uzyskanie optymalnego ustawienia naziemnej anteny Starlink.

Anteny Starlink są samonastawne, poruszane za pomocą silnika. Po ustawieniu anteny we właściwej pozycji, nie musi się ona więcej poruszać. Dzieje się tak, ponieważ oprócz silników używanych podczas wstępnego pozycjonowania, antena może elektronicznie sterować swoją wiązką za pomocą układu przesuwników fazowych.

Nowe wersje anten Starlink przeznaczone do montażu na dachach kamperów lub samolotów nie są w ogóle sterowane mechanicznie a jedynie elektronicznie.

Hakowanie anten

Anteny Starlink nie są przeznaczone do demontażu przez użytkowników. Próba demontażu może spowodować uszkodzenie i utratę gwarancji, ale niektórzy hakerzy i tak to robią.

Wiele osób zdemontowało swoje anteny, aby zobaczyć, co jest w środku lub aby przystosować stacjonarne anteny do użytku mobilnego w samochodach, lub pieszego. Chociaż obecnie dostępne są anteny mobilne, nie zawsze tak było.

Demontaż anteny

Istnieje niewiele oficjalnych informacji na temat budowy anten stacji naziemnych Starlink. Zasadniczo wiemy tylko tyle, ile odkryli hakerzy (rysunki 10, 11 i 12).

Antena Starlink jest niezwykle skomplikowanym urządzeniem i stanowi najbardziej krytyczną część sprzętu naziemnego. Po obejrzeniu filmów z prób rozbioru anten można stwierdzić, że jest to arcydzieło inżynierii. Antena zawiera wiele układów elektronicznych, w tym procesor oparty na architekturze ARM, pamięć RAM i wiele niestandardowych układów scalonych. Prawdopodobnie wszystkie te elementy służą do sterowania układem przesuwników fazowych.

Do filmów ilustrujących demontaż anten należą:

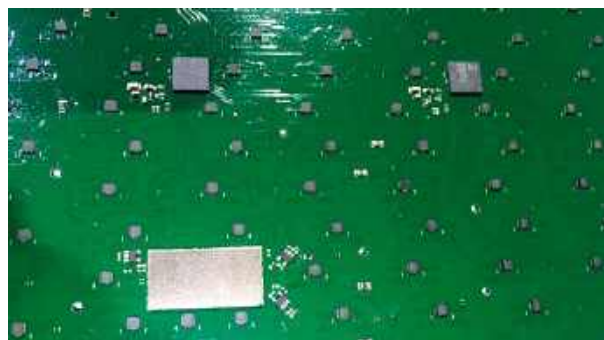
- Starlink Teardown: DISHY DESTROYED! <https://youtu.be/iOmdQnIlmRo>
- TSP #181 – Starlink Dish Phased Array Design, architektura



Rysunek 8. standardowa antena Starlink dla zwykłych użytkowników domowych (po lewej), z polem widzenia 100°. Antena o wysokiej wydajności (środkowa) jest przeznaczona dla firm i przedsiębiorstw, ponieważ może łączyć się z większą liczbą satelitów, jest bardziej odporna na ekstremalne warunki eksploatacji i ma pole widzenia 140°. Płaska antena o wysokiej wydajności (po prawej) jest przeznaczona do zastosowań mobilnych, takich jak samochody kempingowe i łodzie, również z polem widzenia 140°. Źródło: Starlink



Rysunek 9. prostokątna antena Starlink skierowana do góry jest widoczna z przodu samolotu. Źródło: Starlink



Rysunki 10 i 11. Część płytki PCB anteny Starlink. Ścieżki PCB są zakrzywione, aby zapewnić stałą długość wszystkich ścieżek (a sygnały RF nie lubią ostrych zakrętów). Źródło: <https://youtu.be/AlvIWF0AXIO>

i dogłębna analiza RF, <https://youtu.be/h6MfM8EFkGg>

- Starlink Dish TEARDOWN! – część 1 – SpaceX BugBounty jest otwarte podczas publicznej bety Starlink, <https://youtu.be/QuDtSo5tpLk>
- Starlink Dish TEARDOWN! – część 2 – Konsola szeregową i monit logowania. Czy odgadniesz hasło Dishy'ego? https://youtu.be/38_KTq8j0Nw
- Starlink RECTANGLE Teardown Details – Praca nad uproszczeniem anteny Rectangle w celu stworzenia panelu o niskim poborze mocy, <https://youtu.be/AlvIWF0AXIO>

Dobry artykuł na ten temat można znaleźć na stronie siliconchip.au/link/abjf

Telefonia komórkowa

W sierpniu 2022 roku firma Starlink nawiązała współpracę z operatorem T-mobile w Stanach Zjednoczonych, w celu świadczenia usług telefonii komórkowej za pośrednictwem



Rysunek 12. Część stosu układu fazowego po nieskomponowanej stronie płytki drukowanej anteny. Źródło: <https://youtu.be/AlvIWF0AXIO>

satelitów V2 Starlink. Pierwsze testy miały miejsce w 2023 roku. W przeciwieństwie do innych systemów telefonii satelitarnej, ten nowy będzie wykorzystywał standardowe urządzenia mobilne. Początkowo usługa pozwoli wysłać wiadomości tekstowe i nawiązywać połączenia głosowe.

Całkowita przepustowość każdego z satelitów będzie wynosić 2...4 Mb/s, co odpowiada 1000...2000 jednoczesnym połączeniom głosowym lub milionom wiadomości tekstowych. Intencją jest wykorzystanie tej usługi na odludnych obszarach poza zasięgiem sieci komórkowych lub w sytuacjach awaryjnych. Początkowo usługa będzie oferowana tylko w USA, ale T-mobile będzie współpracować z dostawcami Internetu w innych krajach.

Wyzwania technologiczne związane z zapewnieniem łączności satelitarnej dla

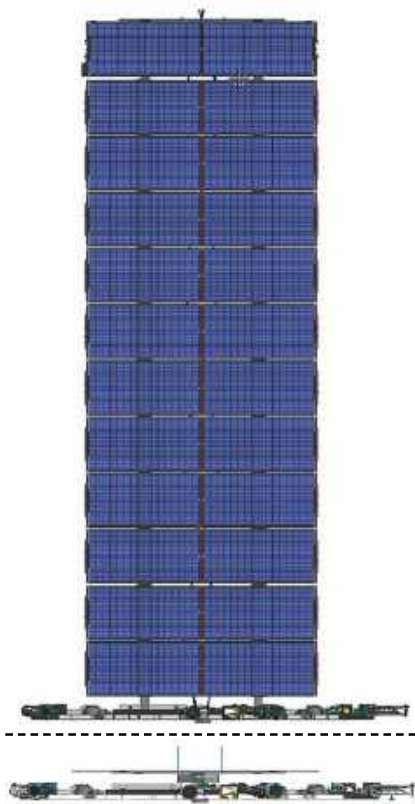
standardowego telefonu komórkowego są poważne. Po pierwsze, zanim sygnał telefoniczny pokona około 550 lub więcej kilometrów, jego amplituda ulegnie silnemu osłabieniu. Ponieważ satelita porusza się z prędkością około 27 000 km/h, należy wziąć pod uwagę znaczne przesunięcie częstotliwości, wynikające z efektu Dopplera.

Telefon będzie elektronicznie śledzony za pomocą anteny sterowanej fazowo, która w trakcie ruchu satelity po swojej orbicie może kierować wiązkę w jego kierunku. Według Elona Muska są to najbardziej zaawansowane technicznie anteny sterowane fazowo na świecie.

Satelity wykorzystywane w tej usłudze będą bardzo duże, o długości 7 m i masie 1,25 tony każdy, a antena będzie miała wymiary 5×5 m i będzie składana na czas startu. Takie anteny

Tabela 3. Proponowane powłoki orbitalne i liczba satelitów Starlink V2 (łącznie 29 988)

Nachylenie względem orbity	Wysokość orbity	Liczba płaszczyzn orbitalnych	Liczba satelitów/płaszczyznę	Łączna liczba satelitów
53,0°	340 km	48	110	5280
46,0°	345 km	48	110	5280
38,0°	350 km	48	110	5280
96,9°	360 km	30	120	3600
53,0°	525 km	28	120	3360
43,0°	530 km	28	120	3360
33,0°	535 km	28	120	3360
148,0°	604 km	12	12	144
115,7°	614 km	18	18	324

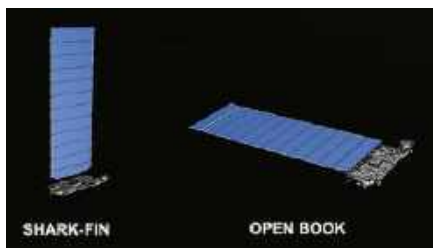


Rysunek 13. Satelity Starlink mogą obniżyć swój profil, aby uniknąć kolizji. Źródło: <https://astronomy.com/news/2022/02/spacex-defends-starlink-over-collision-concerns>

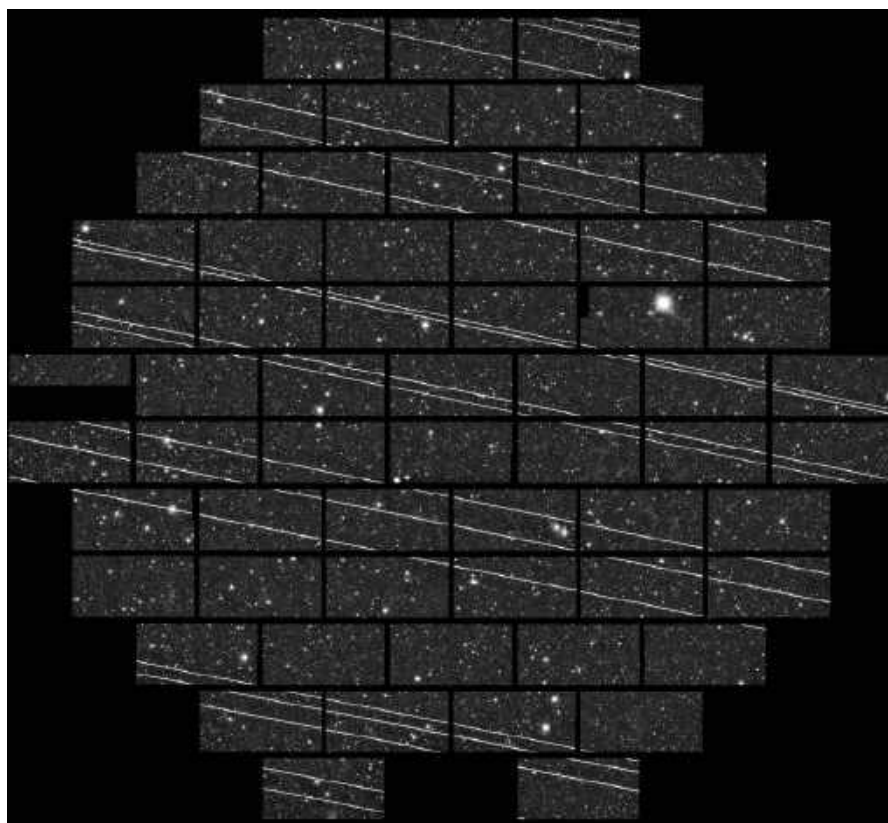
są zbyt duże dla rakiety SpaceX Falcon 9, więc zostaną wyrzucane rakiety SpaceX Starship. Firma SpaceX zaproponowała również miniaturową wersję satelity V2, która zmieści się w rakiecie Falcon 9.

Każdy satelita V2 będzie obsługiwał jedną komórkę telefonii komórkowej o powierzchni prawie 17 000 km². Docelowo będzie pracowało 30 000 satelitów V2 (tabela 3), co wystarczy do pokrycia całej powierzchni Ziemi, wynoszącej około 510 milionów km²!

Dopóki cała konstelacja satelitów V2 nie zostanie uruchomiona, łączność z telefonem komórkowym będzie możliwa tylko wtedy, gdy satelity V2 będą widoczne dla danego użytkownika.



Rysunek 15. konfiguracja płetwy rekina zmniejsza ilość światła słonecznego odbijanego w kierunku Ziemi. Źródło: <https://astronomynow.com/2020/05/05/spacex-to-debut-satellite-dimming-sunshade-on-next-starlink-launch/>



Rysunek 14. Zdjęcie z 2019 r. wykonane w Międzyamerykańskim Obserwatorium Cerro Tololo (CTIO) w Chile po wystrojeniu drugiej partii satelitów Starlink. Ta 333-sekundowa ekspozycja zawiera 19 smug z satelitów. Źródło: <https://noirlab.edu/public/images/iotw1946a/>

Również samochody Tesli będą mogły łączyć się z siecią komórkową Starlink na obszarach objętych zasięgiem sieci T-mobile lub innych obszarach obsługiwanych przez innych dostawców.

Oprócz usługi komórkowej, satelity V2 będą również zapewniać łączność internetową za pośrednictwem konwencjonalnych stacji naziemnych lub powietrznych.

Unikanie kolizji i żywotność satelitów

Satelity Starlink, a właściwie wszystkie obecnie używane satelity muszą mieć możliwość wykonywania manewrów, aby uniknąć kolizji i utrzymywać się na swojej orbicie. Muszą również mieć możliwość zejścia z orbity pod koniec swojego życia, aby zapobiec gromadzeniu się nadmiernej ilości śmieci wokół Ziemi.

Satelity Starlink są wyposażone w silniki Halla (HET), elektryczne silniki jonowe, które wykorzystują krypton jako materiał napędowy do wykonywania wymaganych manewrów. Nawet jeśli pod koniec życia satelity pędnik ulegnie awarii, zejdzie on z orbity w ciągu około czterech lat, wejdzie w atmosferę Ziemi i spłonie.

Unikanie kolizji z dużą liczbą satelitów znajdujących się obecnie w przestrzeni

kosmicznej ma kluczowe znaczenie dla uniknięcia syndromu Kesslera. Jest to zjawisko, w którym kolizja satelitów generuje dużą ilość śmieci. Odłamki te powodują kolejne kolizje i powstają nowe odłamki. Prowadzi to do efektu kaskadowego i sprawia, że przestrzeń orbitalna może stać się bezużyteczna.

W systemie Starlink jest wykorzystany autonomiczny system unikania kolizji oparty na sztucznej inteligencji, pobierający dane do śledzenia z systemu osiemnastej eskadry obrony kosmicznej US Space Force.

Przypuśćmy, że satelita Starlink zbliży się do innego obiektu i nie będzie mógł zejść mu z drogi. W takim przypadku może on złożyć swój panel słoneczny, aby mieć mniejszą powierzchnię, a tym samym zmniejszyć ryzyko kolizji (rysunek 13).

Utracone satelity Starlink

Satelity Starlink są początkowo rozmieszczane na znacznie niższych orbitach niż te, na których później działają. Ma to na celu przeprowadzenie wstępnych testów technicznych. Jeśli satelita okaże się niefunkcyjnym, jego orbita szybko się zacieśni i satelita spłonie w atmosferze ziemskiej. To zapobiega gromadzeniu się śmieci orbitalnych. Jeśli satelita przejdzie

testy prawidłowo, jego orbita zostanie podniesiona.

4 lutego 2022 49 satelitów V1.5 (Grupa 4-7) zostało umieszczonych na niskiej orbicie testowej. W tym czasie wystąpiła poważna burza geomagnetyczna. Spowodowało to zwiększony opór w ruchu, a 38 satelitów spłonęło w atmosferze Ziemi. Pozostałych 11 zostało przesuniętych na wyższą orbitę.

Wpływ satelitów na badania astronomiczne

Od samego początku realizacji projektu Starlink, przewidującego wystrzelenie tysięcy satelitów, astronomowie mieli obawy dotyczące utrudnień w prowadzeniu swoich obserwacji nieba. Na **rysunku 14** przedstawiony jest jeden z pierwszych przypadków zakłóceń obrazu spowodowanych wystrzeleniem partii satelitów Starlink w listopadzie 2019 roku.

Do strategii łagodzących wspomniane zakłócenia należą:

- Użycie osłony o nazwie VisorSat zakrywającej anteny radiowe i inne części satelity. Jest ona przezroczysta dla fal radiowych, ale eliminuje odbicia światła (**rysunek 17**).
- Użycie powłoki pochłaniającej o nazwie Dark Sat, jednak sprawia ona, że satelity zbyt szybko się nagrzewają, więc preferowane jest użycie wspomnianej osłony VisorSat.
- Konfiguracja panelu słonecznego w formie otwartej księgi, o wysokim współczynniku odbicia, która tuż po starcie zmienia się na konfigurację w formie płetwy rekina, z panelem skierowanym w stronę Ziemi (**rysunek 15**).
- Testowane są również manewry przechylania satelitów podczas podnoszenia orbity, aby zminimalizować odbicia światła (**rysunek 16**).

Hakowanie systemu Starlink

W sposób naturalny system Starlink stał się celem hakerów. Nie zalecamy kontynuować tej działalności, ale przedstawiamy uzyskane informacje w formie ciekawostki. Na stronie <https://github.com/KULeuven-COSIC/Starlink-FI> grupa hackerów opublikowała materiał „Glitched on Earth by Humans: A Black-Box Security Evaluation of the SpaceX Starlink User Terminal”, wyjaśniający sposób wykonania dowolnego kodu na terminalu użytkownika Starlink (**rysunek 18**).

Nie ma to żadnego określonego celu poza eksperymentowaniem. Spodziewamy się, że dotychczas wykryte luki w zabezpieczeniach systemu zostały już dawno załatwane.

ZMIANY ORIENTACJI SATELITY USTAWIENIE MATRYCY PODCZAS PODNOSZENIA ORBITY

Obracający się satelita sprawia, że światło słoneczne odbija się od węższej krawędzi matrycy, zmniejszając intensywność odbić



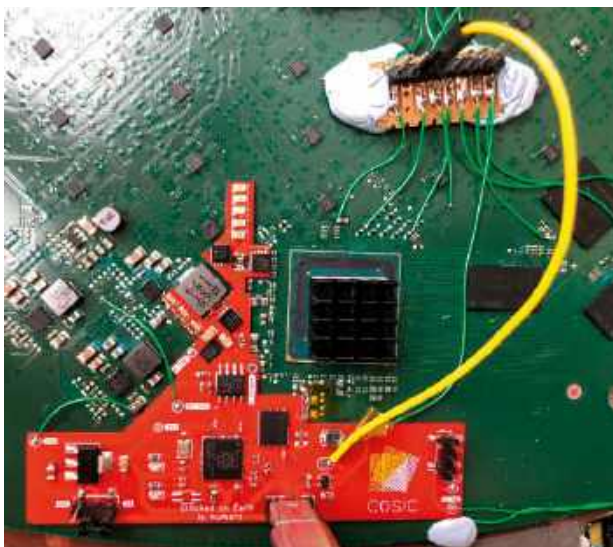
Rysunek 16. Szczegóły konfiguracji płetwy rekina. Źródło: takie samo jak na rysunku 14

VISORSAT OGRANICZANIE POWIERZCHNI ANTEN

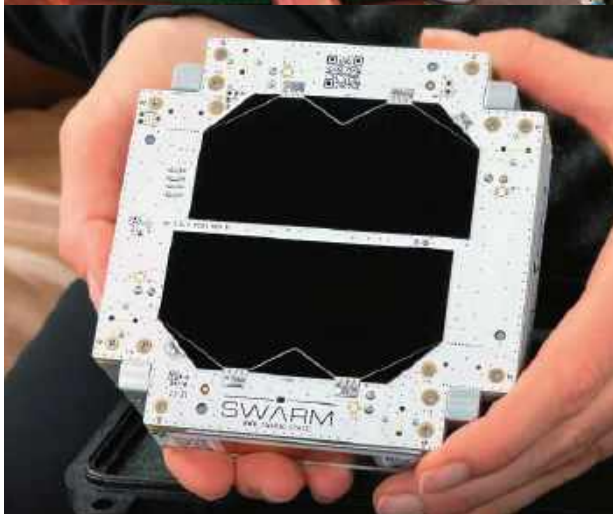
Ostona przeciwsłoneczna blokuje dostęp światła słonecznego do anten, zapobiegając odbiciom



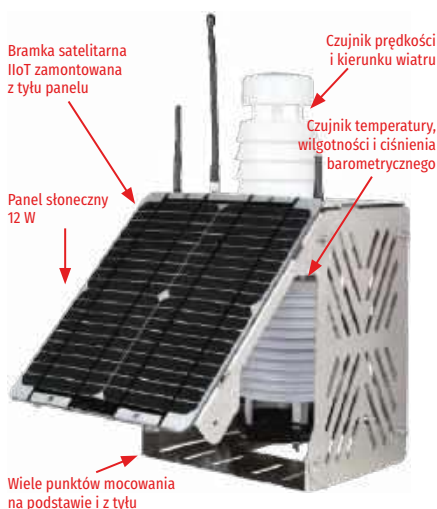
Rysunek 17. Ostona została dodana do późniejszych satelitów Starlink, aby zmniejszyć ilość światła odbijanego w kierunku Ziemi. Źródło: takie samo jak na rysunku 15



Rysunek 18. Płytki „Modchip” (czerwona) i interfejs dodany do panelu anteny Starlink. Źródło: <https://github.com/KULeuven-COSIC/Starlink-FI>



Rysunek 19. Satelita Swarm SpaceBEE, najmniejszy satelita komercyjny



Rysunek 20. Przykład komercyjnej, zdalnej stacji pogodowej ModuSense z wbudowanym systemem łączności Swarm. Źródło: www.freewave.com/products/modusense-weather-station/



Rysunek 21. Urządzenie śledzące Swarm dopasowane do potrzeb. Źródło: <https://swarm.space/swarm-announces-new-asset-tracking-product/>

Ta działalność hackerska nie przeszkadza administratorom systemu Starlink, którzy w ramach akcji „Bug Bounty Program” zachęcają do dalszych prób złamania systemu. Spółka Starlink zapłaci 25 000 USD każdemu, kto znajdzie błąd w ich sieci. Osoby chętne do sprawdzenia swoich sił powinny odwiedzić stronę www.siliconchip.au/link/abjh

Ponadto grupa entuzjastów z Uniwersytetu Teksaskiego w Austin opracowała sposób wykorzystania sygnałów Starlink jako alternatywy dla systemu GPS. Informacje na ten temat są dostępne na stronie www.siliconchip.au/link/abji

Spółka Swarm

Spółka Swarm (<https://swarm.space/>) oferuje globalną łączność IoT (Internet of Things) o niskiej przepustowości, realizowaną za pośrednictwem satelitów SpaceBEE (**rysunek 19**) – BEE oznacza „podstawowe elementy elektroniczne”.

Swarm Technologies stała się spółką zależną SpaceX w lipcu 2021 roku. Co ciekawe,

oddział venture capital amerykańskiej agencji wywiadowczej CIA o nazwie In-Q-Tel, wymienia system Swarm jako jeden ze swoich start-up’ów (<https://www.iqt.org/portfolio/>).

Satelity wykorzystywane przez spółkę Swarm są uważane za najmniejsze komercyjnie aktywne satelity, o rozmiarach ¼U (11×11×2,8 cm) i masie około 400 g. ¼U to oznaczenie Cubesat odnoszące się do rozmiarów satelity w stosunku do standardowego sześciangu 1U o wymiarach 10×10×10 cm, chociaż ściśle rzecz biorąc, satelity Swarm nieznacznie wykraczają poza standard Cubesat.

Satelity Swarm są klasyfikowane jako pikosatelity. Znajdują się one na orbicie synchronicznej ze Słońcem na wysokości 450...550 km, a ich planowana konstelacja liczy 150 obiektów.

Orbita synchroniczna ze Słońcem to specjalny rodzaj orbity polarnej, przebiegającej (mniej więcej) z północy na południe, na której satelita odwiedza to samo miejsce na powierzchni Ziemi o tej samej porze każdego dnia. Na stronie <https://kube.tools.swarm.space/pass-checker/> można sprawdzić, kiedy następny satelita Swarm znajdzie się w danej okolicy.

Panele słoneczne i baterie służą do zasilania

urządzeń SpaceBEE, a antena rozkłada się, gdy satelita znajdzie się we właściwym miejscu na orbicie.

Wszystkie działające satelity muszą być śledzone w celu uniknięcia wzajemnych kolizji oraz do dokonywania korekt położenia na orbicie. Istniały obawy co do możliwości śledzenia tych satelitów ze względu na ich niewielkie rozmiary, ale i ten problem został rozwiązany.

Satelity są wyposażone w pasywnie odbijające siłę odbitego sygnału radarowego.

Każdy z satelitów jest wyposażony w odbiornik GPS i wysyła dane lokalizacyjne na każde żądanie.

Obecność anteny o długości 1 m poprawia widoczność satelity dla naziemnych radarów śledzących i innych czujników (np. amerykańskiej sieci nadzoru kosmicznego).

Jedną z głównych zalet systemu Swarm, oprócz jego globalnej dostępności, jest niski koszt użytkowania. Urządzenia Swarm i abonamenty transmisji danych są łatwo dostępne dla hobbystów, są również wykorzystywane przez profesjonalnych użytkowników.

Zgodnie z informacjami ze strony internetowej Swarm, typowy abonament transmisji danych dla jednego urządzenia kosztuje 5 USD miesięcznie pozwala na transmisję 750 pakietów danych (do 192 bajtów na pakiet lub 144 kB miesięcznie), w tym do 60 pakietów danych typu downlink (dwukierunkowych), zapewnia szyfrowanie AES256-GCM dla bezpiecznej transmisji, roczną umowę bez konieczności rekonfiguracji sprzętu lub ukrytych opłat oraz dostarczanie danych za pośrednictwem interfejsu API REST lub Webhook do dowolnej usługi w chmurze.



Rysunek 22. Płytkę modemu SparkFun M138 to urządzenie w środku z napisem „Swarm”. Źródło: www.electronics-lab.com/sparkfun-swarm-m138-modem-satellite-transceiver-breakout-board/



Rysunek 23. Zestaw czujników umieszczony na szczycie góry, połączony z systemem Swarm. Źródło: <https://swarm.space/>

Ta ilość danych powinna być wystarczająca dla codziennych odczytów ze zdalnej stacji pogodowej, takiej jak ten pokazany na **rysunku 20**.

Do urządzeń służących do połączenia z systemem Swarm należą tracker zasobów za 99 do globalnego śledzenia zasobów z jedną akwizycją GPS co dwie godziny, z jedną transmisją w dwugodzinny oknie i z wykrywaniem ruchu. Szybkość transmisji danych wynosi 1 kb/s, a wykorzystywane częstotliwości mieszczą się w zakresach od 137 MHz do 138 MHz (downlink) i od 148 MHz do 150 MHz (uplink).

Urządzenie waży 227 g, a bateria wystarcza na ponad 40 dni pracy. Urządzenie może być podłączone do zasilacza zewnętrznego. Dostęp do danych można uzyskać z poziomu Swarm Hive (**rysunek 21**).

Innym urządzeniem Swarm jest modem M138, przeznaczony do wbudowania w sprzęt IoT innej firmy, z danymi dostarczany za pośrednictwem interfejsu API REST lub Webhook do dowolnej usługi w chmurze. Te moduły mają cenę 89 USD przy minimalnym zakupie 25 sztuk.

W przypadku mniejszej liczby urządzeń na stronie www.sparkfun.com/products/19236 można nabyć płytkę SparkFun M138 Modem Breakout Board za 149,95 lub jej nowszą wersję, za 199,95 USD (**rysunek 22**).

Modem M138 ma postać karty Mini PCB Express o wadze 9,6 g i zawiera odbiornik GNSS dla GPS i innych systemów nawigacyjnych. Dane są wysyłane do modemu jako szesnastkowy ciąg ASCII, a dwuliterowe polecenia podobne do NMEA (National Marine Electronics Association) są wysyłane przez

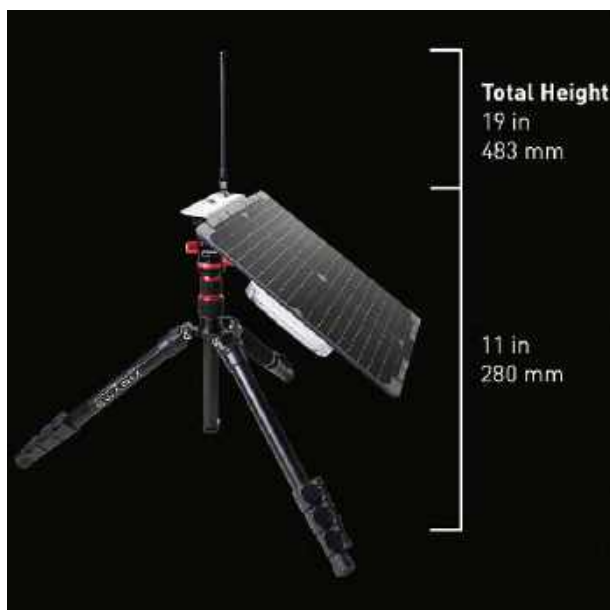
łącze szeregowe 3,3 V (UART). Modem M138 jest wbudowany we wspomniany powyżej moduł śledzenia zasobów.

Aplikacje dla modemu M138 z płytką Breakout pozwalają na odczyt danych ze zdalnych czujników, takich jak stacje pogodowe, urządzenia do zdalnego monitorowania sprzętu, śledzenia zasobów i monitorowania środowiska (**rysunek 23**).

Na koniec, zestaw Swarm Eval Kit za 449 USD (**rysunek 24**) został stworzony, aby dostarczyć deweloperom łatwą w użyciu platformę, z dołączoną płytką FeatherS2 – ESP32 + OLED, portem USB-C i portem I²C dla czujników. Dodatkowe moduły FeatherWing stwarzają wiele dodatkowych możliwości konstrukcyjnych.

W skład zestawu Eval Kit wchodzi stacytów, panel słoneczny, baterie oraz zintegrowana antena VHF i GPS. Bieżący odczyt szumu z tła radiowego pomaga osiągnąć możliwie najlepszą jakość połączenia. Urządzenia można podłączyć za pośrednictwem interfejsu Wi-Fi (tryb AP lub STA), USB lub łączy szeregowych, a danymi można zarządzać za pośrednictwem aplikacji Swarm Cloud i REST API.

Szybkość transmisji danych wynosi 1 kb/s przy maksymalnym rozmiarze pakietu 192 bajtów i zapewnione jest szyfrowanie metodą AES256 GCM. Format poleceń to dwuliterowa NMEA. Zestaw jest dostarczany ze wspomnianym powyżej modemem M138 i waży 2,6 kg.



Rysunek 24. Zestaw Swarm Eval Kit. Dokumentację można znaleźć na stronach <https://swarm.space/documentation-swarm/> i www.sparkfun.com/products/19236 w zakładce „documents”

Uwagi dotyczące dokładności i aktualności opisu

Dołożyliśmy wszelkich starań, aby zapewnić najbardziej dokładne i aktualne informacje, na temat konkretnych szczegółów satelitów Starlink i ich liczby na orbicie. Takie informacje nie są publikowane lub mogą ulec zmianie, ponieważ plany komercyjne SpaceX zmieniają się z czasem.

Pamiętaj, że Starlink, Swarm i Starshield to systemy, które są budowane nawet w chwili, gdy to czytasz, a plany stale ewoluują.

System Starshield

Starshield (www.spacex.com/starshield/) jest pochodną systemu Starlink przeznaczoną specjalnie dla amerykańskiego rządu i wojska. Według danych ze strony internetowej SpaceX, początkowo wykorzystanie systemu Starshield koncentrowało się na obserwacji Ziemi, komunikacji i hostingu transmisji danych.

Usługa związana z obserwacją Ziemi polega na wystrzeliwaniu satelitów z urządzeniami pomiarowymi i dostarczaniu przetworzonych danych bezpośrednio do użytkownika końcowego (agencji rządowej). Do tego dochodzi globalna komunikacja za pomocą sprzętu Starshield, który ma jeszcze wyższy poziom bezpieczeństwa niż Starlink, gdyż transmisja jest szyfrowana metodą end-to-end.

W ramach hostingu transmisji danych budowane są odpowiednie magistrale satelitarne, dostosowane do potrzeb klienta. Magistrala satelitarna to podstawowy element strukturalny statku kosmicznego z wyposażeniem zapewniającym transmisję poleceń i danych, w tym łączność, zasilanie, napęd, kontrola termiczna, kontrola położenia i naprowadzanie.

Istnieje możliwość zainstalowania specjalistycznego wyposażenia wymaganego przez klienta, takiego jak zestaw czujników dostosowanych do konkretnej misji. Jest to tańsze niż budowa od podstaw osobnego satelity. Magistrala opiera swoje działanie istniejących satelitach Starlink V1.5 i V2.0, których panele słoneczne mają dużą powierzchnię.

W razie potrzeby satelity Starshield mogą współdziałać z satelitami Starlink za pośrednictwem systemu międzysatelitarnej komunikacji laserowej.

Aplikacje Starlink mogą być szybko rozwijane dzięki dostępności nośników SpaceX do wynoszenia satelitów na orbitę, dzięki produkcji satelitów i zdolności do szybkiego rozmieszczania ich dużej liczby podczas jednego startu.

Podobne systemy satelitarne

AST SpaceMobile ast-science.com

Firma AST uruchamia niskoorbitálną, szerokopasmową usługę komórkową, która pozwoli na korzystanie ze standardowych, niezmodyfikowanych smartfonów za pośrednictwem satelity z dużą anteną sterowaną

fazowo, o powierzchni 64,4 m². Jego prototyp BlueWalker 3 wystrzelony w listopadzie 2022 roku krąży na orbicie o wysokości 508...527 km i ma pole widzenia równe 777 tys. km².

Firma AST SpaceMobile planuje docelowo rozmieścić konstelację 243 satelitów BlueBird na orbitach o wysokości 725...740 km pod koniec 2023 roku. Satelity BlueBird są podobne do prototypu BlueWalker 3. Późniejsze wersje będą miały jeszcze większy system antenowy. Partnerami AST są firmy AT&T, Vodafone, Orange i Rakuten Mobile.

BlueWalker 3 został wystrzelony w ramach transportu współdzielonego na rakiecie SpaceX Falcon 9 wraz z innymi satelitami Starlink.

Globalstar www.globalstar.com/en-ap

Firma Globalstar oferuje konstelację satelitów na orbicie LEO, umieszczonych na wysokości 1400 km, służących do głosowych połączeń telefonicznych i transmisji danych o niskiej prędkości, za pomocą specjalnych telefonów. W konstelacji znajdują się 24 satelity drugiej generacji. W USA i Kanadzie użytkownicy iPhone'a 14 mogą wysyłać wiadomości alarmowe za pośrednictwem tego samego systemu satelitarnego.

Hughes Network Systems hughes.com

Firma Hughes Network Systems jest amerykańskim dostawcą globalnych, szerokopasmowych usług internetowych, głównie w oddalonych obszarach na Ziemi. Oferuje również komórkowe usługi dosyłowe (backhaul) za pośrednictwem satelitów geostacjonarnych oraz usługi internetowe w samolotach.

Komórkowe usługi backhaul są świadczone za pośrednictwem satelitów, ponieważ transmisja przewodowa lub tradycyjne łącza mikrofalowe prowadzące do odległych lokalizacji są zbyt drogie. Ponieważ wykorzystywane są satelity geostacjonarne, występuje problem dużych opóźnień. To oznacza, że system nie nadaje się do prowadzenia wideokonferencji i do gier, a w komunikacji głosowej występuje znaczne opóźnienie.

Inmarsat www.inmarsat.com

Firma Inmarsat wykorzystuje 14 satelitów umieszczonych na orbicie geostacjonarnej

i oferuje szereg usług, w tym łączność dla 160 000 statków i 17 000 samolotów, a także agencji rządowych i dużych firm.

Do oferowanych usług należą: śledzenie obiektów, dostawa szybkiego Internetu, usługi w stanie zagrożenia i usługi dla zachowania bezpieczeństwa. Do połączenia z systemem Inmarsat wymagane jest użycie specjalnego telefonu lub innego urządzenia końcowego. Samolot Malaysia Airlines, który w tajemniczy sposób zaginął, korzystał z usług telefonii satelitarnej Inmarsat, a analiza danych wykazała, że w ostatniej chwili przed zniknięciem leciał nad południową częścią Oceanu Indyjskiego.

Iridium www.iridium.com

System Iridium wykorzystuje 66 aktywnych satelitów, umieszczonych na polarnej orbicie okołoziemskiej o 100-minutowym okresie obiegu, w sześciu płaszczyznach orbitalnych, oddalonych od siebie o 30°, na wysokości 780 km. Komunikacja odbywa się za pośrednictwem sprzętu Iridium lub urządzeń pochodzących od innych dostawców (www.iridium.com/products/). Usługa zapewnia transmisję danych tekstowych i cyfrowych, odbiór komunikatów SOS, wiadomości głosowych i innych.

Wykorzystywane częstotliwości mieszczą się w przedziale od 1616,0 MHz do 1626,5 MHz, zaś bramka uplink od 29,1 GHz do 29,3 GHz, bramka downlink od 19,1 GHz do 19,6 GHz, a łącza międzysatelitarne od 22,55 GHz do 23,55 GHz.

Kuiper Systems LLC

www.aboutamazon.com/news/tag/project-kuiper

Firma Kuiper Systems jest spółką zależną koncernu Amazon. Jej celem jest zapewnienie ogólnie dostępnego i niedrogiego, szerokopasmowego Internetu satelitarnego społecznościom na całym świecie.

Kuiper buduje konstelację 3276 satelitów na orbicie LEO, a prototypowe satelity mają zostać wystrzelone na początku 2023 roku. Satelity będą orbitować na wysokości 590...630 km.

Lynk Global – lynk.world

Firma Lynk chce stworzyć rodzaj komórkowej stacji BTS umieszczonej w kosmosie,

aby standardowe telefony komórkowe mogły łączyć się z satelitami rozmieszczonymi na orbicie LEO, na wysokości 500 km. Celem jest zapewnienie komunikacji komórkowej w krajach trzeciego świata, aby tamtejsi użytkownicy mogli korzystać z przystępnych cenowo telefonów. Komunikacja będzie także możliwa na obszarach, w których nie ma dostępu do sygnału komórkowego lub dostęp został przerwany z powodu klęski żywiołowej.

Projekt firmy Lynk jest obecnie w fazie testowej i będzie wymagał użycia 1000 satelitów do pełnego pokrycia szerokopasmowego, które spodziewa się osiągnąć do 2025 roku. Docelowa pełna konstelacja będzie zawierać 5000 satelitów.

O3b www.ses.com

System O3b wykorzystuje konstelację 20 satelitów umieszczonych na średniej orbicie okołozemskiej (MEO), na wysokości 8000 km, co pozwala na uzyskanie relatywnie niskich opóźnień w transmisji. Celem jest zapewnienie połączeń internetowych w obszarach wiejskich na szerokościach geograficznych od 50°N do 50°S, obejmujących 96% populacji Ziemi, dla operatorów sieci komórkowych, operatorów telekomunikacyjnych, przedsiębiorstw i instytucji rządowych.

Przykładowo, chodzi o telemedycynę, bankowość elektroniczną i wirtualne nauczanie w takich miejscach jak Samoa, Brazylia, Czad, Timor Wschodni, Papua i Nowa Gwinea. Usługi telefonii komórkowej 4G+ mogą być oferowane w miejscach takich jak Wyspy Cooka poprzez świadczenie usług typu backhaul. Mogą również zapewniać komunikację internetową na morzu, na przykład na statkach wycieczkowych.

W przypadku usług nowej generacji, firma O3b uruchamia satelity typu mPower dla rządu, wojska i różnych przedsiębiorstw i będzie dysponować 11 satelitami na orbicie MEO, z których każdy może wytworzyć 5000 cyfrowo uformowanych wiązek skierowanych do różnych użytkowników.

OneWeb oneweb.net

Firma OneWeb jest w trakcie wystrzeliwania konstelacji 648 satelitów na orbitę LEO, na wysokość 1200 km, w celu świadczenia globalnych usług szerokopasmowego dostępu do Internetu do końca 2023 roku. Klientami mają być instytucje rządowe, wojskowe, firmy telekomunikacyjne i odległe społeczności, a nie osoby prywatne.

Orbcomm – orbcomm.com/pl/partners/connectivity/satellite

FIRMA Orbcomm oferuje konstelację satelitów Isat Data Pro na orbicie GEO i ORBCOMM OG2 oraz LEO, do satelitarnej łączności IoT (Internet of Things).

Project Loon to nieistniejąca już propozycja wykorzystania balonów umieszczanych na dużych wysokościach (18...25 km) do stworzenia rozległej sieci bezprzewodowej. Manewrowanie w celu utrzymania stałej pozycji polegało na dostosowaniu wyporności balonu w celu znalezienia wiatru we właściwym kierunku.

Linki

- Obserwuj gołym okiem satelity Starlink. Informacje na stronie <https://findstarlink.com/>. Należy pamiętać, że satelity Starlink są obecnie mniej widoczne niż kiedyś ze względu na środki

podjęte w celu zminimalizowania zakłóceń w pracy astronomów. Dostępna jest również aplikacja na urządzenia z systemem Android i iOS.

- Sprawdź aktualną lokalizację konstelacji Starlink, a także konstelacji OneWeb i GPS. Informacje na stronie <https://satellitemap.space/>.
- Na stronie www.starlink.com/map dostępna jest interaktywna mapa umożliwiająca określenie dostępności usługi satelitarnej w danej lokalizacji. Chociaż dostępność usługi jest globalna, nadal konieczne jest zawarcie umów krajowych i ustaleń dotyczących rozliczeń.
- Na stronie <https://youtu.be/d29jURzZGe0> dostępny jest film zatytułowany „The Truth About Starlink RV! Is It Worth It?”, zawierający ocenę działania systemu Starlink podczas podróży kamperem po Australii.
- Film przedstawiający kolejną konstelację satelitów krótko po wystrzeleniu, przed ich umieszczeniem na właściwych orbitach, zatytułowany „Starlink Satellites train seen in the sky” jest dostępny na stronie <https://youtu.be/ihVuz8uM1qU>.
- Bardzo ciekawy i prosty projekt pozwalający na odbiór sygnałów nawigacyjnych Starlink za pomocą komputera Raspberry Pi, odbiornika SDR i anteny satelitarnej jest opisany na stronie siliconchip.au/link/abjm. ■

dr David Maddison

Artykuł reprodukowano na podstawie umowy z magazynem „Silicon Chip”, 2022. www.siliconchip.com.au

REKLAMA

Mnóstwo doskonałych projektów, tylko na:

EP.com.pl



Chirurgia obwodowa

Zniekształcenia i obwody zniekształcające, część 3

W poprzednich dwóch odcinkach tego cyklu zajmowaliśmy się zniekształceniami – efektem nieliniowości układów (na przykład wzmacniaczy), która wpływa na kształt przebiegów na wyjściu. Zniekształcenia są zwykle cechą niepożądaną, a projektanci układów starają się je minimalizować. Wielkość niepożądanych zniekształceń mierzy się zazwyczaj współczynnikiem całkowitych zniekształceń harmonicznych (THD), szczególnie w zastosowaniach audio. Współczynnik ten, wraz z podstawowymi pojęciami dotyczącymi zniekształceń, został omówiony w pierwszym artykule tego cyklu („Practical Electronics”, czerwiec 2022 r.; EdW 1/2026).

W poprzednim artykule badaliśmy widma sygnałów w kontekście zniekształceń. Widmo jest to rozkład natężenia sygnału w funkcji częstotliwości, ukazujący jakie częstotliwości składowe występują w sygnale. Każdy przebieg okresowy można utworzyć poprzez zsumowanie pewnej ilości fal sinusoidalnych o różnych częstotliwościach i amplitudach. Ta suma znana jest jako szereg Fouriera. Wykres widma sygnału okresowego jest zbiorem prążków położonych przy określonych częstotliwościach.

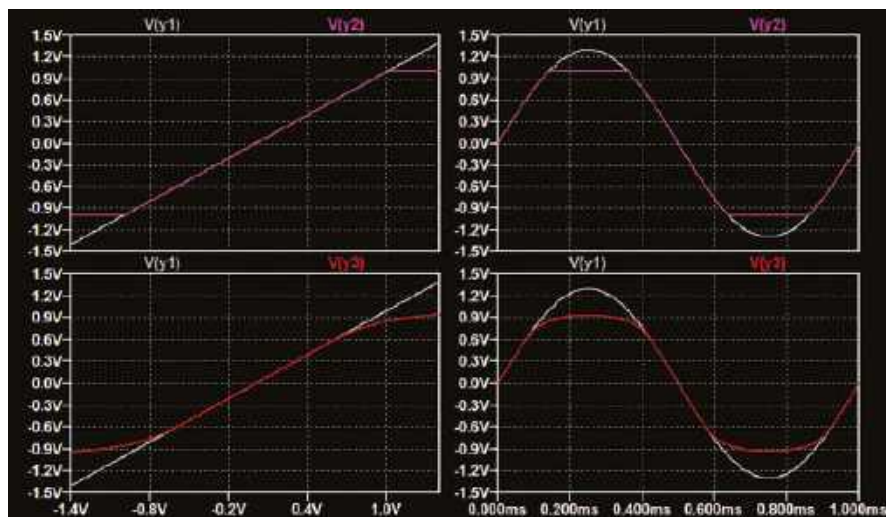
Widmo na wyjściu układu liniowego nie będzie zawierało częstotliwości, które nie występowały na wejściu – chociaż względne amplitudy oryginalnych częstotliwości mogą ulec zmianie. Jeśli z powodu nieliniowości układu występują zniekształcenia sygnału, na wyjściu pojawią się częstotliwości, których na wejściu nie było. W przypadku sinusoidalnego sygnału wejściowego, dodatkowe częstotliwości na wyjściu będą wielokrotnościami częstotliwości wejściowej, czyli harmonicznymi sygnału wejściowego. Jest to podstawa ilościowego określenia zniekształceń poprzez współczynnik THD. Wykreślenie widma sygnału i obliczenie THD jest możliwe w symulatorze LTSpice. Użycie LTSpice zapewnia wszechstronny wgląd w zagadnienie zniekształceń. Posługiwanie się symulatorem wymaga jednak pewnej uwagi i ostrożności, co było w zeszłym odcinku kluczowym elementem dyskusji.

Zniekształcenia są na ogół niepożądane. Mają one jednak swoje zastosowania, między innymi w efektach dźwiękowych wykorzystywanych przez muzyków. Chyba najbardziej znanym przykładem są podłogowe efekty zniekształcające używane przez gitarzystów elektrycznych. Niniejsza seria artykułów została w istocie zainspirowana układami do efektów muzycznych autorstwa Johna Clarke’a, zamieszczanymi w „Practical Electronics” w ciągu ostatnich dwóch lat. W tym miesiącu przyjrzymy się układom, które

można wykorzystać do zamierzonego wytwarzania zniekształceń.

Obcinanie

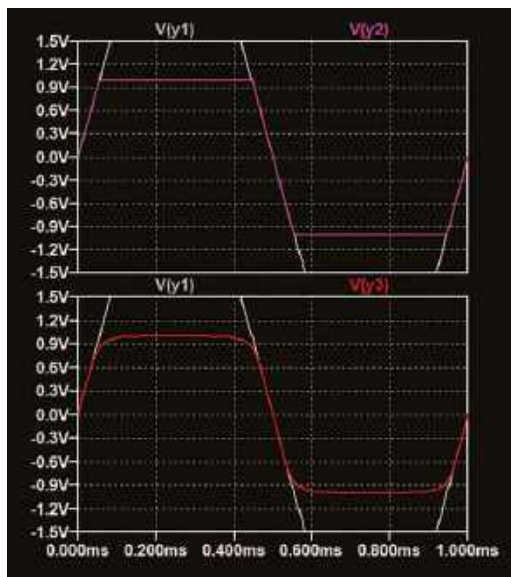
Większość zniekształceń w zastosowaniach muzycznych uzyskuje się z pomocą układów, które powodują obcinanie sygnału. Amplituda przebiegu zostaje tak ograniczona, że jego szczyty ulegają spłaszczeniu. Efekt ten bywa również nazywany „nasyceniem”. Zanim jednak przejdziemy do omówienia układów, które można wykorzystać do uzyskania takiego efektu, podsumujmy, co rozumiemy



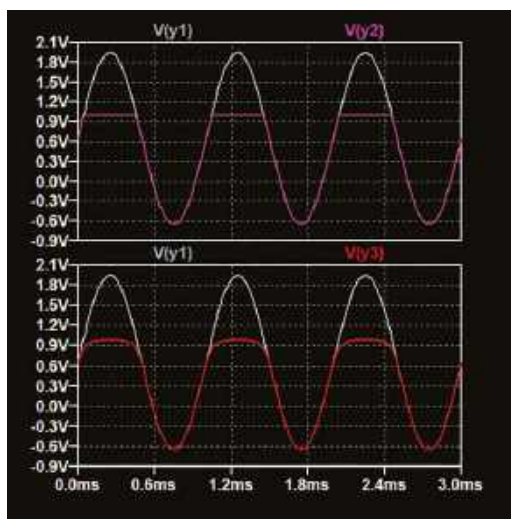
Rysunek 1. Obcinanie (clipping) – twarde (na górze) i miękkie (na dole)

przez obcinanie, i przyjrzyjmy się kilku zasadniczym jego odmianom („twarde”/„miękkie” i symetryczne/asymetryczne).

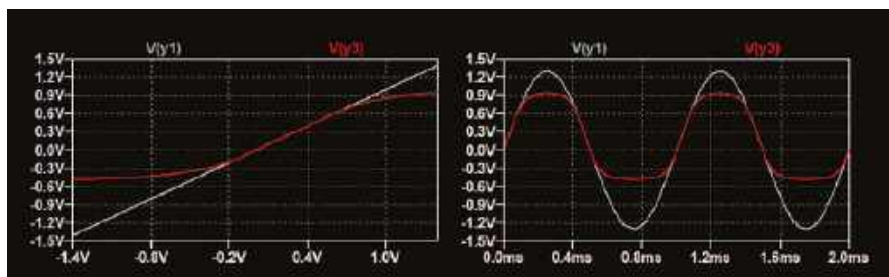
Jeśli dla idealnego wzmacniacza wykreśliśmy funkcję przenoszenia (napięcie wyjściowe w funkcji napięcia wejściowego), będzie



Rysunek 2. Ta sama funkcja obcinająca/ograniczająca (clipping), ale przy większym sygnale wejściowym



Rysunek 3. Zniekształcenia asymetryczne – uzyskane poprzez przepuszczenie fali sinusoidalnej ze składową stałą przez funkcję przenoszenia przedstawione po lewej stronie rysunku 1



Rysunek 4. Zniekształcenia asymetryczne wynikające z asymetrycznej funkcji przenoszenia

ona dla wszystkich możliwych napięć idealną linią prostą (szara linia $V(y1)$ na wykresach po lewej stronie **rysunku 1**). Wzmacniacze rzeczywiste cechuje jednak określona maksymalna amplituda wyjściowa, co powoduje, że ich funkcje przenoszenia są podobne do tych pokazanych przez kolorowe krzywe $V(y2)$ i $V(y3)$ po lewej stronie **rysunku 1**. Obie krzywe różnią się między sobą ostrością przejścia od obszaru liniowego (dla małych amplitud) do obszaru pełnego ograniczania (dla dużych amplitud). Przechodzenie stosunkowo szybkie (jak krzywa $y2$) nazywane jest „obcinaniem twarde” (ang. „hard clipping”; [przypis redaktora](#)). Do przechodzenia bardziej stopniowego (jak krzywa $y3$) odnosi się określenie „obcinanie miękkie” (ang. „soft clipping”; [przypis redaktora](#)). Wykresy na **rysunku 1** uzyskano w oparciu o matematycznie zdefiniowaną funkcję przenoszenia, omówioną w pierwszym artykule tego cyklu. Szare krzywe po prawej stronie **rysunku** ukazują nieznkształconą falę sinusoidalną na wyjściu.

W przypadku konkretnego wzmacniacza – lub innego układu, który wprowadza obcinanie – stopień zniekształceń zależy od amplitudy sygnału wejściowego. Ilustruje to **rysunek 2**, na którym widać przebieg wyjściowy układu o tej samej funkcji przenoszenia, co na **rysunku 1**, ale przy ponad dwukrotnie większej wartości szczytowej na wejściu (1,3 V na **rysunku 1**, 3 V na **rysunku 2**). Wraz ze wzrostem amplitudy sinusoidalnego sygnału wejściowego, przebieg na wyjściu układu przesterowującego ma tendencję do stawania się falą prostokątną. Łagodniejsza funkcja przesterowania sprawia, że narożniki fali prostokątnej są zaokrąglone. Taki sam efekt da filtrowanie sygnału zniekształconego w filtrze dolnoprzepustowym.

Funkcja przenoszenia układu zniekształcającego nie musi być

symetryczna. Jeśli będzie ona asymetryczna, wystąpi niesymetryczne zniekształcenie. Ten sam efekt da wysterowanie układu o symetrycznej funkcji przenoszenia sygnałem z nałożoną składową stałą. Przypadek taki ilustruje **rysunek 3**. Mamy tu taką samą amplitudę sygnału, co na **rysunku 1** (1,3 V w szczycie), ale wejściowa fala sinusoidalna ma teraz nałożoną składową stałą o wartości +0,65 V. Ujemna wartość szczytowa wynosi -1,0 V i tylko nieznacznie przekracza napięcie graniczne funkcji obcinającej. Ujemna połówka sygnału jest więc zniekształcona tylko w niewielkim stopniu. Silnie obciążona jest natomiast dodatnia wartość szczytowa. Składową stałą sygnału z **rysunku 3** można łatwo usunąć, używając kondensatora sprzęgającego.

Rysunek 4 przedstawia asymetryczną funkcję przenoszenia z obcinaniem „miękkim” i wynikowy przebieg wyjściowy dla fali sinusoidalnej 1,3 V bez składowej stałej. Jak podawaliśmy w pierwszym artykule cyklu, funkcja obcinania „miękkiego” wykorzystuje źródło behawioralne o funkcji:

$$\text{sgn}(v(x)) * \text{upLim}(\text{abs}(v(x)), \text{lim}, 0, 4)$$

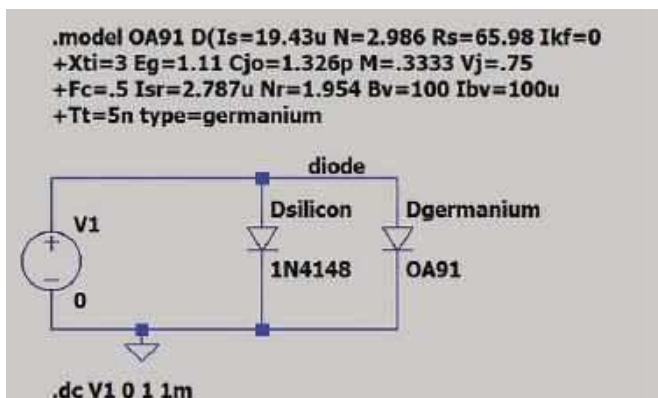
W tym przypadku lim jest wartością graniczną czyli napięciem przycinania. W przypadku przycinania asymetrycznego użyto dwóch wersji takich funkcji z wartościami granicznymi odpowiednio 0,5 i 1,0. Funkcje te są wybierane przez funkcję „if”:

$$V = \text{if}(v(x) + 0,5, a, b)$$

a i b to dwie funkcje ograniczające. W symulatorze LTspice funkcja $\text{if}(x, y, z)$ zwraca y , jeśli $x > 0,5$, w przeciwnym razie zwraca z . Schemat symulowanego układu LTspice można znaleźć w pierwszym artykule tego cyklu.

Terminologia

Do opisu efektów zniekształcających stosowanych przez muzyków są często używane określenia „booster” („podbijacz”), „overdrive” („przesterowanie”; żargonowo „przester”), „distortion” („zniekształcenie”) czy „fuzz”. Termin „overdrive” pierwotnie odnosił się do zwiększania głośności wzmacniacza gitarowego na tyle, by spowodować jego przesterowanie i uzyskać w ten sposób dźwięk pełniejszy i bardziej agresywny. W początkowej erze gitar elektrycznych, przed powszechnym wprowadzeniem efektów gitarowych, efekt zniekształcania dźwięku trzeba było uzyskiwać w samym wzmacniaczu. Wczesne wzmacniacze gitarowe były jeszcze budowane nie na tranzystorach, lecz na lampach elektronowych, a lampa, ze względu

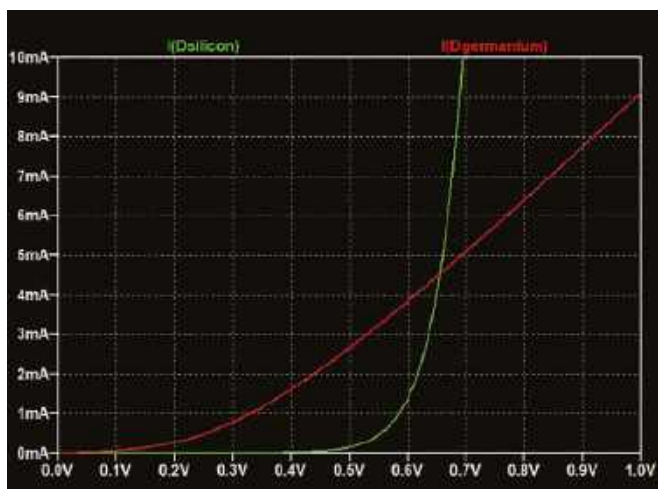


Rysunek 5. Układ w LTspice służący do uzyskania charakterystyki prądowo-napięciowej diody

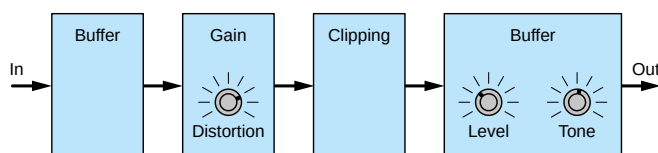
na swoją charakterystykę, dawała przesterowanie stosunkowo łagodne. Dzisiejsze efekty podłogowe reklamowane jako „overdrive” zapewniają zazwyczaj łagodny rodzaj przesterowania i mają na celu właśnie symulację brzmienia przesterowanego wzmacniacza lampowego.

Efekt „booster” realizuje po prostu wzmocnienie dźwięku bez przesterowania. Stosowane są ewentualnie filtry zapewniające różne wzmocnienia dla różnych częstotliwości, co zmienia barwę dźwięku. Ale, jak wspominaliśmy wyżej i jak widać na rysunku 2, zwiększenie amplitudy sygnału powoduje większe zniekształcenia. Wzmocnienie sygnału może wywołać przesterowanie dalszych stopni wzmacniacza lub zwiększyć zniekształcenia w układzie przesterowującym (urządzeniu efektywnym) podłączonym do wyjścia „boostera”.

Efekty „overdrive”, „distortion” i „fuzz” są ściśle związane z przesterowywaniem sygnału. „Overdrive”, jak już nadmienialiśmy, z reguły odnosi się do przesterowania łagodniejszego. Efekty „distortion” i „fuzz” zapewniają zazwyczaj „twarde” przesterowanie, ewentualnie występują obwód „miękkiego” przesterowania sygnałem o stosunkowo dużej amplitudzie.



Rysunek 6. Charakterystyka prąd/napięcie diody krzemowej (zielona) i germanowej (czerwona)

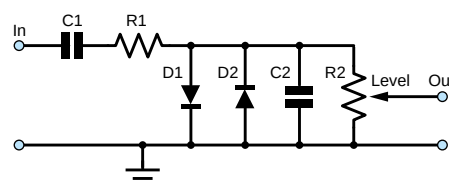


Rysunek 7. Schemat blokowy urządzeń do efektów zniekształceniowych

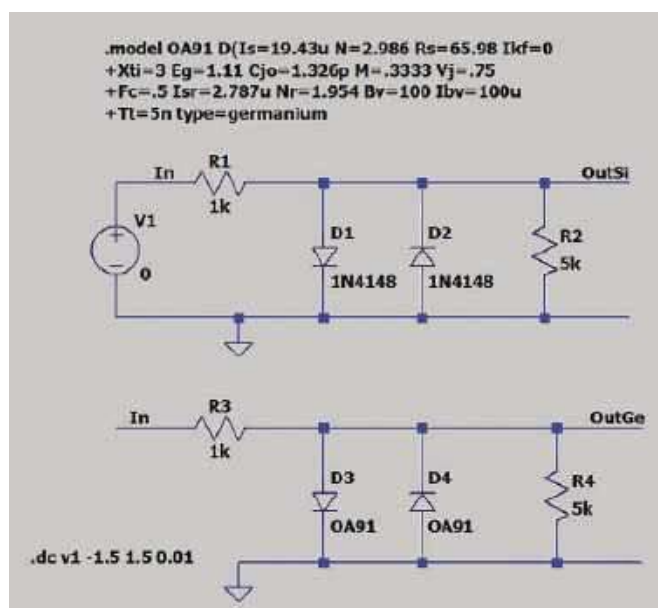
Oba efekty, a szczególnie „fuzz”, dają sygnał o kształcie bardzo zbliżonym do prostokątnego. Przypis redaktora: stopniem końcowym układu „fuzz” jest często cyfrowy przerzutnik Schmitta, wytwarzający niemal idealny przebieg prostokątny. W rozbudowanych „fuzzach” wyjście przerzutnika może jeszcze taktować ciąg przerzutników bistabilnych, dzielących częstotliwość sygnału przez 2, 4, ..., a więc dających dodatkowe dźwięki niższe o oktawę, dwie itd.

Nieliniowe funkcje przenoszenia układów ograniczających (na przykład lewa strona na rysunkach 1 i 4), oprócz tego, że wprowadzają zniekształcenia, wpływają też na dynamikę (zakres głośności) sygnału wyjściowego. „Miękką”, zakrzywiającą się charakterystyką powoduje, że mniejsze sygnały są wzmacniane bardziej niż większe. Proces ten znany jest jako kompresja. W przypadku gitary powoduje to podtrzymanie wybrzmiewania granych dźwięków. Drgania szarpniętej struny instrumentu z biegiem czasu słabną, a głośność dźwięku się zmniejsza. Jeśli wraz ze spadkiem sygnału wzrasta wzmocnienie, głośność będzie miała tendencję do pozostawania na bardziej wyrównanym poziomie niż sygnał pochodzący bezpośrednio z instrumentu, co spowoduje dłuższe trwanie dźwięków. Odnosi się to do sytuacji przedstawionej na rysunku 1, zakładając, że fala sinusoidalna na rysunku ma taką amplitudę, jak maksymalna amplituda dźwięku. W tym przykładzie zniekształcanie są tylko dźwięki najgłośniejsze. Z kolei sytuacja na rysunku 2 odpowiada przypadkowi, że nawet dość ciche dźwięki powodują mocne przesterowanie. Wtedy zmiany głośności wyjściowej przy różnych poziomach na wejściu będą stosunkowo niewielkie.

W poprzednim artykule tego cyklu szczegółowo omówiono fakt, że zniekształcenia zmieniają skład częstotliwościowy sygnału poprzez dodanie do niego harmonicznich, a w przypadku fal o kształcie innym niż sinusoidalne – również innych częstotliwości. Aby z efektu zniekształcającego uzyskać jak najlepszy dźwięk, celowe może być odfiltrowanie



Rysunek 8. Układ ograniczający z diodami



Rysunek 9. Układ w LTspice do symulowania diodowego obwodu ograniczającego, z różnymi rodzajami diod

wyższych częstotliwości filtrem dolnoprzepustowym. A stosując takie czy inne filtry na wejściu, dźwięk można zniekształcać w różnym stopniu w różnych zakresach częstotliwości. Wiele efektów zniekształcających zawiera, oprócz regulacji stopnia zniekształcenia, różnego rodzaju regulatory barwy tonu (filtry), którymi można nastawiać brzmienie dźwięku.

Charakterystyka diody

Układy obcinające są zasadniczo ogranicznikami napięcia. Naturalnym ogranicznikiem napięcia może być dioda, która jest znana z tego, że spadek napięcia na niej jest niemal stały w szerokim zakresie prądu. Zamiast zmuszania całego wzmacniacza do nasycania się, możemy uzupełniać go o obwody zniekształcające (ograniczające) oparte na diodach. Ograniczenie diodowe jest podstawą wielu, choć nie wszystkich, urządzeń do uzyskiwania efektów zniekształcających.

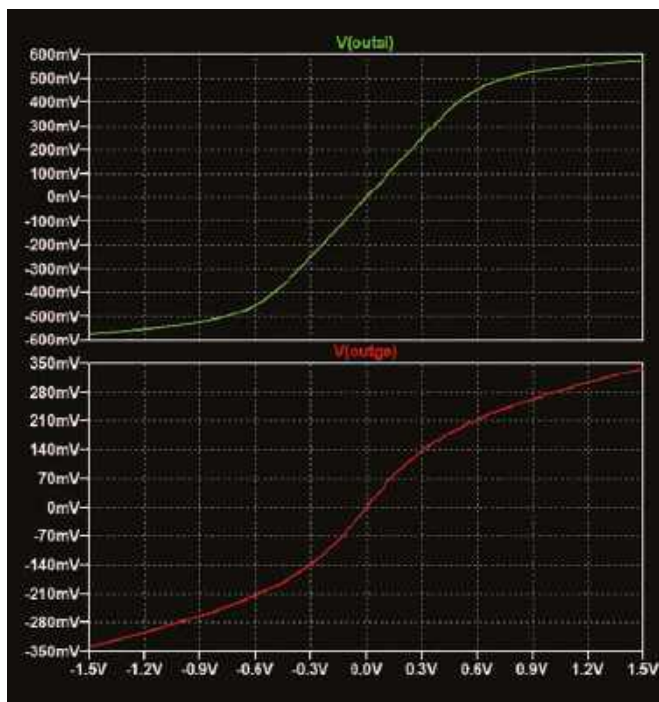
Możemy wykreślić charakterystykę prądu w funkcji napięcia (I/V) diody w LTspice, używając stałoprądowej symulacji przemiatającej – na przykład w układzie pokazanym na **rysunku 5**. Badamy tutaj charakterystyki I/V diody krzemowej 1N4148, bardzo często stosowanej, a także diody germanowej OA91. Diody germanowe były niegdyś szeroko stosowane, zanim w większości zastosowań nie zostały zastąpione przez elementy krzemowe. Diody krzemowe mają znacznie lepsze parametry: prąd upływu, maksymalne napięcie wsteczne, stabilność, maksymalną temperaturę roboczą i cenę. Diody germanowe znajdują jednak nadal pewne niszowe zastosowania, a jednym z nich są właśnie układy zniekształcające. Typ OA91 został w tym przypadku wybrany nieco arbitralnie, ponieważ został znaleziony zaraz po rozpoczęciu poszukiwań w Internecie. Wśród modeli dostarczanych wraz z symulatorem LTspice nie ma żadnych modeli diod germanowych. Na rysunku 5 pokazano odpowiednią instrukcję .model.

Rysunek 6 przedstawia wyniki symulacji układu z **rysunku 5**. Widać dwie zasadnicze cechy. Po pierwsze, dioda germanowa zaczyna przewodzić prąd przy niższym napięciu niż dioda krzemowa (około 0,2...0,3 V zamiast 0,6...0,7 V). Po drugie, dioda germanowa charakteryzuje się mniej gwałtownym wzrostem prądu przy wzroście napięcia przewodzenia – co przekłada się na łagodniejszą odpowiedź układu ograniczającego, w którym używana jest dioda.

Układy do efektów zniekształcających

Rysunek 7 przedstawia schemat blokowy urządzenia wytwarzającego efekty zniekształcające. Jest to bardziej ilustracja ogólnej koncepcji toru sygnałowego niż struktura jakiegoś rzeczywistego układu, ponieważ w konkretnych implementacjach funkcje bloków mogą być połączone. Często stosuje się przełącznik obejściowy, nie pokazany tutaj, który może przekierować sygnał poza cały układ i umożliwić muzykowi ominięcie efektu. Bardziej złożone układy tego typu mogą mieć rozbudowane tor sygnałowy, na przykład umożliwiające wybór jednego z kilku efektów lub ich miksowanie.

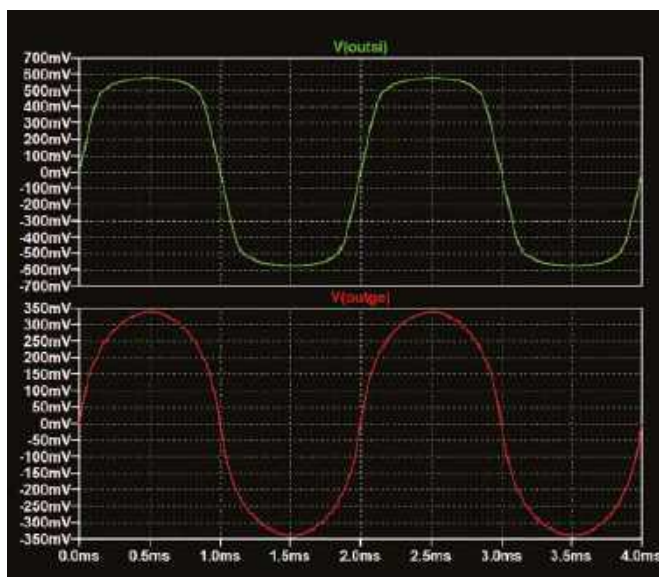
Układ na **rysunku 7** zaczyna się od przedwzmacniacza-bufora, który zapewnia odpowiednie parametry wejściowe dla źródła sygnału (na przykład przetwornika gitary elektrycznej). Bufor może zawierać filtry, na przykład w celu usunięcia zakłóceń radiowych o wielkiej częstotliwości. Następny stopień to wzmacniacz o regulowanym wzmacnieniu, któryysterowuje obwód przesterowujący. Jak już omawialiśmy i demonstrowaliśmy na **rysunku 1** i **rysunku 2**, amplituda sygnału wchodzącego do obwodu przesterowującego wpływa na stopień wytwarzanych zniekształceń. Zmiana wzmacnienia powoduje zatem zmianę stopnia zniekształcenia przy danym poziomie sygnału wejściowego. Stopień przesterowujący jest obwodem nieliniowym o funkcji przenoszenia z ograniczeniami – jak na **rysunku 1**



Rysunek 10. Wyniki symulacji układu z **rysunku 9**: funkcje przenoszenia obwodów ograniczających z różnymi rodzajami diod (na górze krzemowa, na dole germanowa)

i 4. Charakterystyka stopnia może być „twarda lub „mięka”, symetryczna lub asymetryczna. W układzie mogą być przełączniki wybierające różne warianty układu lub regulatory ciągłe do ustawiania takich parametrów jak na przykład asymetria. Bufor wejściowy i stopień wzmacniający mogą być realizowane przez jeden układ.

Za stopniem przesterowującym mamy już sygnał zniekształcony. Jego amplituda może jeszcze wymagać dostosowania do czułości układu docelowego (wzmacniacza głośnikowego lub następnego efektu w łańcuchu). Zapewnia to widoczny na **rysunku 7** bufor wyjściowy. Niektóre efektowe urządzenia zniekształcające oferują możliwość



Rysunek 11. Wyniki symulacji przejściowej dla sinusoidalnego sygnału na wejściu – diodowy obwód ograniczający z różnymi rodzajami diod (na górze krzemowa, na dole germanowa)

regulacji barwy tonu. Zazwyczaj jest to filtr, który w regulowanym stopniu obcina wysokie częstotliwości. W niektórych konstrukcjach część wyjściowa może nie zawierać aktywnego wzmacniacza, a regulacja poziomu sygnału i barwy tonu może wchodzić w skład obwodu przesterowania.

Ograniczanie na diodach

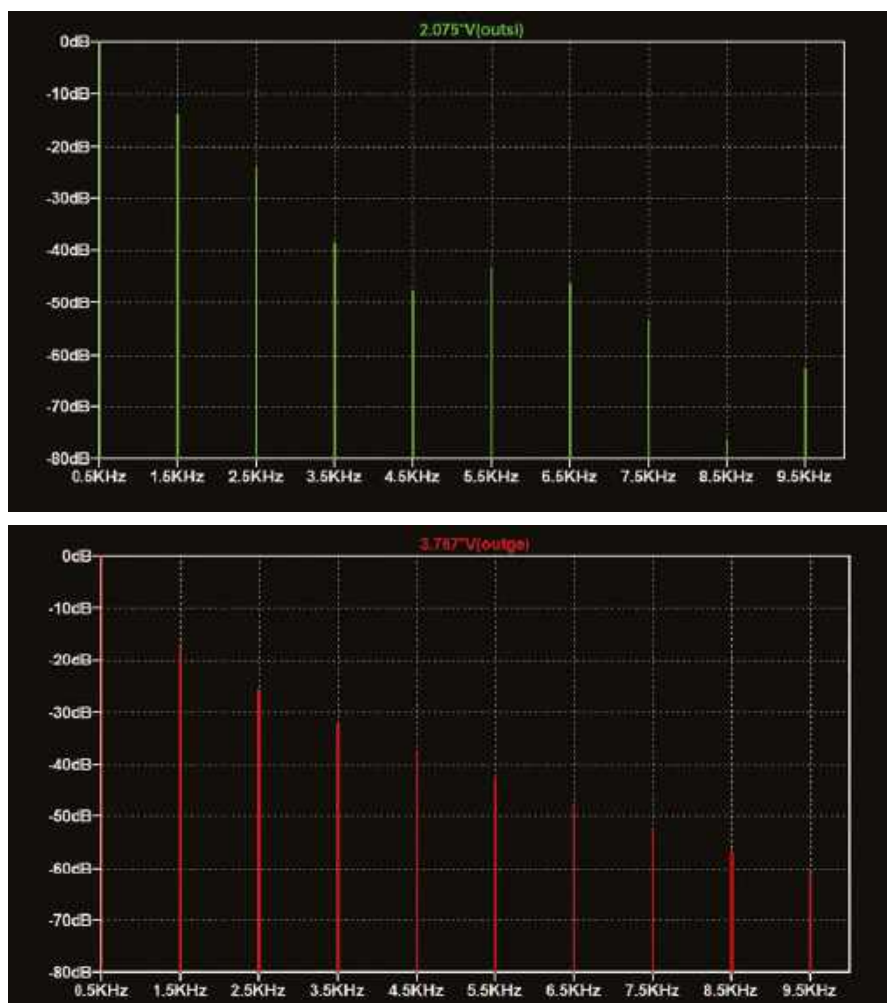
Układ pokazany na **rysunku 8** to typowy obwód przesterowania z diodami stosowany w efektach zniekształcających. Istnieją różne warianty układu i w poszczególnych wariantach nie wszystkie elementy są używane. Układ odbiera zazwyczaj sygnał ze wzmacniacza o zmiennym wzmocnieniu, który określa stopień zniekształcenia przebiegów. Wejście jest sprzężone kondensatorem C1, usuwającym składową stałą, zapobiegającym ewentualnemu wpływaniu diod na pracę wzmacniacza sterującego, a także osłabiającym wzmocnienie dla niskich częstotliwości. C1 ma zazwyczaj pojemność kilku mikrofardów.

Sygnał wyjściowy jest ograniczany (obcinany) napięciami przewodzenia diod D1 i D2. Są użyte dwie diody, aby oddziaływać zarówno na dodatnie jak i ujemne półokresy przebiegu. Można uzyskać obcinanie asymetryczne, stosując jako D1 i D2 różne rodzaje diod lub dając w każdym kierunku inną ich liczbę. Używa się tu standardowych diod krzemowych, diod Schottky'ego, germanowych lub LED-ów. Diody krzemowe i germanowe porównano w układzie z **rysunku 5**. Diody Schottky'ego mają podobny kształt charakterystyki jak diody krzemowe, ale z niższym napięciem przewodzenia – wynoszącym około 0,2 V, więc bliższym diodom germanowym – co daje taki efekt jak przesunięcie charakterystyki „krzemowej” na **rysunku 6** w lewo. Diody LED mają napięcie przewodzenia od 1,2 V do 4 V, w zależności od koloru.

R1 ogranicza prąd diod. R1, R2 i C2 tworzą filtr dolnoprzepustowy, użyty w celu usunięcia wyższych harmonicznych ze zniekształconego przebiegu. Część sygnału z diod jest kierowana do wyjścia przez potencjometr R2, zapewniający regulację poziomu. Sygnał za potencjometrem może bezpośrednio stanowić wyjście układu lubysterowywać dalszy stopień wzmacniający, opcjonalnie wyposażony w regulację barwy tonu. R1 i R2 mają zazwyczaj po kilka kiloomów, a C2 – kilka nanofardów.

Przykład symulacji

Układ na **rysunku 9** jest wersją układu z **rysunku 8** wprowadzoną do symulatora LTspice. Pominięto kondensator filtra dolnoprzepustowego C2. Dzięki temu



Rysunek 12. Widma sygnałów z **rysunku 11** (na górze „krzem”, na dole „german”)

można w pełni obserwować wpływ diod na zawartość harmonicznych. Nie ma również kondensatora sprzęgającego na wejściu (C1) – nie jest on potrzebny, gdyż mamy idealne źródło sygnału (bez składowej stałej), a poza tym sprzężenie stałoprądowe jest niezbędne do obliczania w symulatorze funkcji przenoszenia. Wyniki z **rysunku 10** pokazują, że funkcja przenoszenia przypomina charakterystykę „miękką” z **rysunku 1**, przy czym łagodniejszą charakterystykę ma obwód z diodą germanową.

Jeśli w układzie z **rysunku 9** zmienimy źródło napięcia na generator fali sinusoidalnej i przeprowadzimy symulację przejściową zamiast przemiatania stałoprądowego.

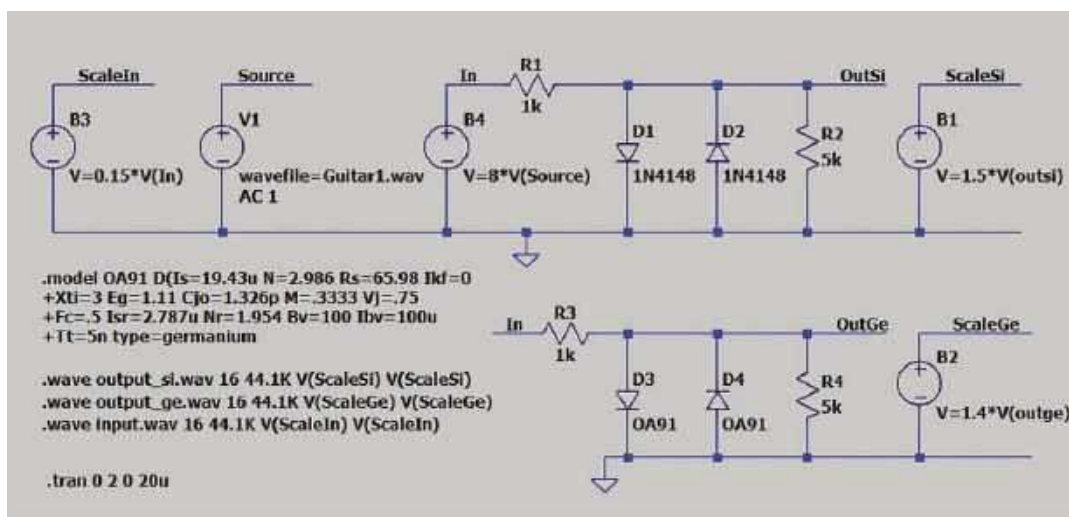
V1 source: SINE(0 1.5 500)

Simulation command: .tran 0 50m 0 10n

Otrzymamy wyniki przedstawione na **rysunku 11**. Ilustrują one fakt, że diody krzemowe powodują silniejsze obcięcie (dają bardziej prostokątny kształt fali) niż diody germanowe.

Wyniki symulacji przejściowej w LTspice można dzięki analizie FFT użyć do uzyskania widma sygnału wyjściowego. Należy wziąć pod uwagę różne czynniki wpływające na skuteczne wykonywanie tej operacji, szczegółowo omówione w zeszłym odcinku tego cyklu. **Rysunek 12** – liniowa skala częstotliwości ułatwia identyfikowanie poszczególnych harmonicznych. Wyniki pokazują, że badane układy wprowadzają harmoniczne nieparzyste, natomiast nie wprowadzają parzystych. Oba typy diod dają różne poziomy względne poszczególnych harmonicznych. Różnica w zawartości harmonicznych spowoduje odmienne efekty w barwie dźwięku – stosowanie układów z różnymi rodzajami diod daje różne brzmienie.

Aby ułatwić porównywanie widm, zostały one znormalizowane tak, aby poziom częstotliwości podstawowej (częstotliwości przebiegu sinusoidalnego na wejściu, 500 Hz) wynosił w obu przypadkach 0 dB. Dokonano tego poprzez pomiar amplitudy prądu o częstotliwości podstawowej w pierwszej wersji symulacji i obliczenie współczynnika



Rysunek 13. Schemat w LTspice do symulacji układów diodowych z wejściem i wyjściem w postaci plików WAV

skalowania wymaganego do przesunięcia tego prążka do poziomu 0 dB. Następnie wykreślono skalowany sygnał. Na przykład początkowo prążek dla częstotliwości podstawowej układu „krzemowego” wynosił $-6,34$ dB. Wymagane skalowanie wyniosło $1/10^{-6,34/20}=2,075$. Wykres został poddany edycji, aby otrzymać wartość $2,075 \cdot V(\text{outs})$. W tym celu należało kliknąć na tytuł wykresu prawym przyciskiem myszy.

Test odstuchowy

Biorąc pod uwagę, że omawiamy tu układy do przetwarzania dźwięku, przydatna byłaby możliwość odsłuchiwania rezultatów. Gdybyśmy faktycznie opracowywali urządzenie do celów muzycznych, ostateczne testy odsłuchowe musielibyśmy przeprowadzać z prawdziwym układem. LTspice zapewnia możliwość posługiwania się sygnałami zapisanymi w formie plików WAV – trochę dla zabawy, ale również w celu szybkiego badania skutków zmian w różnych układach. Można na przykład jako sygnału wejściowego użyć krótkiego nagrania muzycznego, powiedzmy dźwięku gitary, i wyniki symulacji odsłuchiwać z pliku.

Pliki WAV omawialiśmy szczegółowo w wydaniu Practical Electronics z lipca 2020 r., więc tutaj potraktujemy sprawę pobieżnie. W przypadku użycia pliku WAV jako danych wejściowych należy zmienić „wartość” źródła LTspice na postać:

wavefile=filename

gdzie filename to sama nazwa pliku WAV, o ile znajduje się w tym samym folderze co schemat, lub nazwa z pełną ścieżką do pliku, jeśli znajduje się on w innym miejscu. Jeśli plik WAV jest stereo, to domyślnie będzie używany

kanal pierwszy. Aby zapisać wynik do pliku WAV, należy na schemacie umieścić dyrektywę .wave. W przypadku sygnałów audio najlepiej skonfigurować format standardowy, na przykład stereo, 16 bitów, 44,1 kHz (taki jak na płytach kompaktowych). Przykładowo, aby umieścić napięcie out1 w obu kanałach stereo w pliku output1.wav, użyjemy:

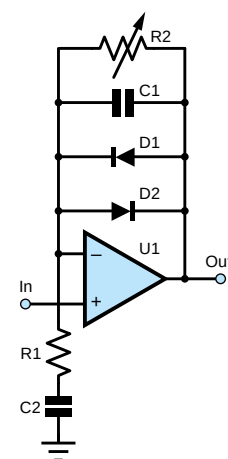
```
.wave output1.wav 16 44,1K V(out1) V(out1)
```

Istotne jest, że maksymalna amplituda sygnału wpisywanego do pliku WAV wynosi 1 V. By dopasować się do tego wymagania, trzeba zazwyczaj skalować zarówno wejście jak i wyjście układu. Po stronie behawioralnych źródeł sygnału jest to proste do wykonania. W przypadku wyjścia jest jednak zwykle konieczne przeprowadzenie symulacji wstępnej, zmierzenie poziomu szczytowego i odpowiednie ustawienie skalowania.

Na **rysunku 13** widzimy wersję układu z rysunku 9, skonfigurowaną dla wejścia i wyjścia w postaci plików WAV. Układ został przetestowany próbką dźwięku gitary ściągniętą z Internetu (freewavesamples.com). Można używać, że każdy z obu rodzajów diod daje inną barwę dźwięku.

Inne układy

Omówiony tutaj obwód z diodami nie jest jedynym układem, jaki można stosować. Częstość rozwiązaniem jest włączenie diod w obwód sprzężenia zwrotnego wzmacniacza operacyjnego. Jest to podejście podobne jak w przypadku wzmacniaczy logarytmicznych, omówionych w „Practical Electronics” w grudniu 2021 r. (EdW 1/2025). Podstawowa, typowa konfiguracja jest pokazana na **rysunku 14**, ale, podobnie jak poprzednio, istnieją różne warianty tego układu.



Rysunek 14. Układ zniekształcający oparty na wzmacniaczu operacyjnym

R1 i R2 wyznaczają wzmocnienie dla sygnałów zmiennych, tak jak w konwencjonalnym wzmacniaczu nieodwracającym. Typowe wzmocnienia wynoszą od 10 do 100. C2 (nie występujący w niektórych wariantach układu) blokuje składową stałą, wskutek czego dla napięć stałych sprzężenie zwrotne wynosi 100% i układ ma jednostkowe wzmocnienie. Diody obcinają sygnał wyjściowy, ograniczając jego amplitudę. Tak jak poprzednio, można użyć różnych rodzajów diod lub różnych ich ilości, uzyskując asymetrię. R2 jest zwykle rezystorem zmiennym – zmienia wzmocnienie układu, co pozwala kontrolować poziom zniekształceń. C1 to kondensator, zazwyczaj niewielki, zmniejszający wzmocnienie układu przy większych częstotliwościach, co odsuwa ryzyko niestabilności. Kształt funkcji przenoszenia tego układu różni się nieco od charakterystyki obwodu z diodowym ogranicznikiem omówionego wcześniej, więc i brzmienie będzie inne. Niektóre podłogowe efekty zniekształcające zawierają oba typy układów.

Inne podejście zastosowano w efekcie Nutube Guitar Overdrive and Distortion Pedal autorstwa Johna Clarke'a („Practical Electronics”, marzec 2021). W układzie tym nie ma diod; do przesterowania jest doprowadzany sam wzmacniacz. Jest to stopień lampowy ze wspólną katodą, oparty na niskonapięciowej triodzie Nutube. Przesterowanie asymetryczne uzyskuje się poprzez polaryzację wejścia napięciem stałym. ■

Ian Bell

Ten artykuł był wcześniej opublikowany na łamach „Practical Electronics”, sierpień 2022 (www.epemag3.com)



Migające diody LED i śliniący się inżynierowie (29)

W chwili, gdy piszę te słowa, siedzę w swoim domowym biurze. Z jakiegoś powodu moja biedna, zagubiona żona (Gina Wspaniała) uparcie nazywa to pomieszczenie jadalnią. Ale to absurd, ponieważ nigdy tu nie jadamy. Nie ma na to miejsca, ponieważ pokój ten jest pełen moich rzeczy.

Och, jakie to przykre!

Wyrzuty sumienia budzi we mnie niestety animatroniczna głowa robota, która leży na stole obok klawiatury i spogląda na mnie ponuro. Potrafię odczytać jej myśli. Brzmią one: „Dlaczego nie piszesz o tym, jaka jestem wspaniała?”.

Jak być może pamiętacie z poprzednich artykułów, mój przyjaciel Steve Manley, który mieszka w Wielkiej Brytanii, oraz ja, obecnie mieszkający w Stanach Zjednoczonych, wspólnie pracowaliśmy nad tym znakomitą artefaktem. Prawdę mówiąc, całą pracę wykonał Steve, ja natomiast spędziłem większość czasu podrygując w tle – jak szalona postać z gry Whac-A-Mole. Podsuwałem tylko przydatne (mam nadzieję) sugestie.

Steve z powodzeniem wykorzystał swój program do modelowania 3D – CAD Fusion 365 (<https://autode.sk/3kMrIA7>) – oraz drukarkę 3D, tworząc dwa wspaniałe egzemplarze głowy, jeden dla siebie, a drugi dla mnie. Zaowocowało to tym, że – jak pisałem w odcinku 27. „Practical Electronics” z maja 2022 r. (EdW 12/2025) – szczęśliwie mogłem trzymać „swoją głowę” w dłoniach, chichocząc do siebie: „ale mam skarb”. A nie jest to coś, co słyszy się codziennie.

Och, co za radość!

Dla waszej przyjemności i rozkoszy Steve z pomocą Fusion 365 stworzył niesamowite, fotorealistyczne renderzy swojej głowy (wiem; gdy to piszę, brzmi to tak samo dziwnie, jak wtedy, kiedy to czytacie). Zaczniemy od dokładnego obejrzenia i przeanalizowania **rysunku 1**, który pokazuje widok głowy od tyłu.

Dwie okrągłe powłoki (u góry) to tyły oczu, z których każde jest wyposażone w jedną z naszych płytek SMAD (Steve and Max's Awesome Display). Każda taka płytka ma 45 trójkolorowych diod LED, więc efekt jest naprawdę spektakularny! Jak niesamowite są te oczy, można zobaczyć na filmie Steve'a na YouTube: <https://bit.ly/3MZ9g3q>.

Ponadto, jak już wcześniej wspominałem, każde oko jest podłączone do dwóch małych

silników serwo (serwomechanizmów), zamontowanych na górnej ramie poziomej. Serwomechanizmy te umożliwiają oddzielne obracanie oczu na boki oraz pochylanie ich w górę i w dół. Steve zaprojektował sprytne rozwiązanie, które pozwala na niezależne wykonywanie obu tych ruchów bez wzajemnego zakłócania się. Nadal wpadam w zachwyt, gdy widzę to w akcji.

Mała płytka drukowana umocowana pośrodku górnej poziomej ramy jest dołączona do modułu głównego poprzez dwuprzewodowy interfejs I²C. Płytkę tę może być używana do sterowania 16 serwomechanizmami (w obecnej implementacji używamy tylko siedmiu). Szczepnie mówiąc, ciągle muszę sobie przypominać, że rysunek 1 to renderowanie, a nie prawdziwe zdjęcie.

Zwróćmy teraz uwagę na dużą niebieską podstawę na dole rysunku. Na pierwszym planie po prawej stronie widzimy większy, mocniejszy serwomechanizm, którego ramię jest przymocowane do dolnej poziomej ramy (będziemy musieli wymyślić jakieś fajniejsze nazwy dla tych elementów konstrukcyjnych). Serwomechanizm ten może być użyty do obracania całej głowy na boki.

Na koniec zwróćcie uwagę na pomarańczowy wspornik w kształcie litery U, który łączy dolną ramę poziomą z poziomą ramą

górną. Sposób połączenia tego wspornika z dolną ramą pozwala na przechylanie głowy na boki (jest to ruch, który w odcinku 25. „Practical Electronics” z marca 2022 r. (EdW 10/2025) określiliśmy jako „przekrzywianie”). Jednocześnie sposób połączenia tego wspornika z górną ramą poziomą pozwala na pochylanie głowy do przodu i do tyłu. Należy pamiętać o tych możliwościach ruchu podczas oglądania **rysunku 2**.

Najważniejsza rzecz, na którą należy zwrócić uwagę, to dwa serwomechanizmy zamontowane na dolnej platformie poziomej, których ramiona są połączone z górną platformą poziomą. W połączeniu z wyżej wymienionym łącznikiem w kształcie litery U i jego mechanizmami mocującymi, napędzanie obu tych serwomechanizmów w górę lub ciągnięcie ich obu w dół powoduje odpowiednio wychylenie głowy w górę lub w dół. Natomiast jeśli jeden serwomechanizm jest napędzany w górę, a drugi pociągany w dół, głowa przechyla się na bok.

Czym właściwie jest serwomechanizm?

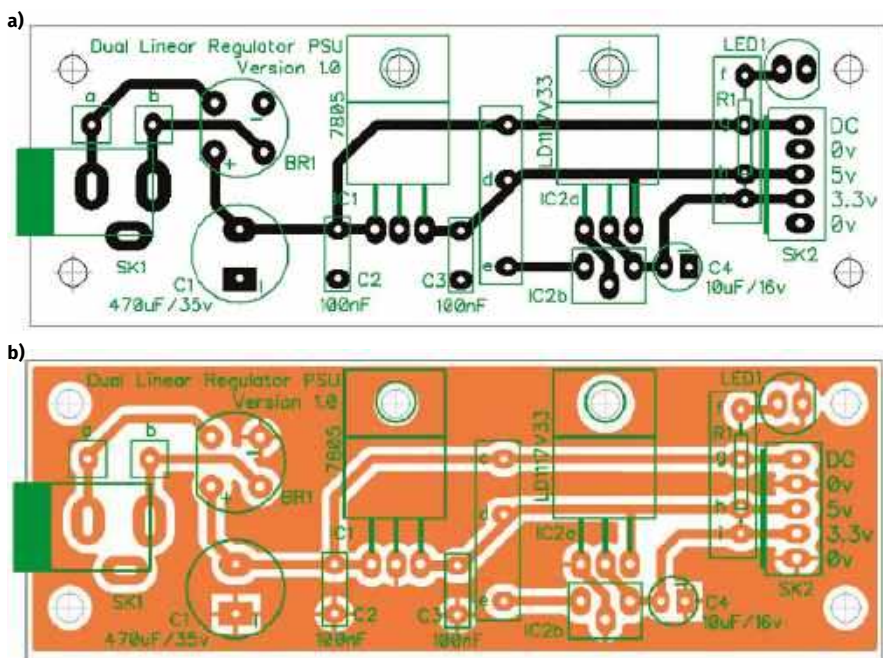
Jest to naprawdę interesujące pytanie. Tak interesujące, że odpowiedź zostawię na następny odcinek. Rzecz w tym, że jest tu wiele do omówienia – w tym różnice między silnikami obrotowymi a silnikami liniowymi, silnikami krokowymi a silnikami serwo, silnikami prądu przemiennego a silnikami prądu stałego... A nawet musimy



Rysunek 1. Widok animatronicznej głowy od tyłu. Zdjęcie: Steve Manley



Rysunek 2. Widok animatronicznej głowy od przodu. Zdjęcie: Steve Manley



Rysunek 3. U góry: a) projekt płytki drukowanej zasilacza Joe Farra tylko ze ścieżkami sygnałowymi; u dołu: b) projekt z dodanym rysunkiem miedzi dla masy

sobie zdefiniować, co rozumiemy pod pojęciem „silnik”. Mam nadzieję, że odpowiedź was zaskoczy.

W międzyczasie Steve podzielił się kilkoma interesującymi informacjami na temat serwo mechanizmów użytych w naszej animatronicznej głowie. Zanim przejdziemy dalej, zauważmy, że kiedy widzimy element określany jako „serwo mechanizm 9 g” lub „serwo mechanizm 55 g”, to nie chodzi tu o to, jaką masę serwo mechanizm może podnieść (co jest typowym błędem początkujących, zwłaszcza że w tym kontekście „podnoszenie” nie jest terminem odpowiednim). Odnosi się to raczej do rzeczywistej masy samego serwo mechanizmu. Jeśli chodzi o to, ile pracy może wykonać serwo mechanizm („praca” ma tu ścisłe znaczenie fizyczne), to w tym miejscu przechodzimy do pojęcia momentu obrotowego, a jest to kolejny złożony temat, o którym porozmawiamy podczas następnego spotkania. Obawiam się jednak, że zbaczymy z tematu – więc wróćmy do Steve’a, który powiedział, co następuje:

„Do funkcji obrotu i pochylenia oczu użyłem czterech serwo mechanizmów 9 g. Ponadto użyłem trzech serwo mechanizmów 55 g do obrotu i pochylenia głowy. W przypadku serwo mechanizmów 9 g zacząłem od Longrunner KY66 ze sklepu Amazon, z plastikowymi przekładniami (<https://amzn.to/3MU6KeI>). Niestety okazało się, że ich jakość była nierówna. Jeden lub dwa działały dobrze, ale pozostałe były albo bardzo głośne, albo działały nierównomiernie.

Jeśli chodzi o serwo mechanizmy 55 g, mój pierwotny wybór padł na elementy Diymor MG996R, również z serwisu Amazon, ale z przekładniami metalowymi (<https://amzn.to/3KXMTU6>). Okazały się one tak samo nierówne jak oryginalne serwo mechanizmy 9 g. Ponadto zakres ruchu dla określonego sygnału wejściowego również był nierówny, co oznaczało, że ruchy pochylenia i przekrzywienia animatronicznej głowy były nierównomierne. W rezultacie nie użyłbym ponownie tych zakupionych na Amazonie serwo mechanizmów i nie polecam ich.

W końcu natrafiłem na firmę HobbyKing z siedzibą w Unii Europejskiej. Firma oferuje szeroką gamę serwo mechanizmów lepszej jakości (<https://bit.ly/37oeJS3>).

Szukałem w Google’u dobrych, uniwersalnych serwo mechanizmów 9 g i znalazłem popularne Hextronic HXT900s (<https://bit.ly/3kN0DNe>). Są to niedrogie serwo mechanizmy o zwartej konstrukcji, z plastikowymi przekładniami. Działają one całkiem dobrze.

Kupiłem również kilka serwo mechanizmów Turnigy TGY-50090 o zbliżonych



Rysunek 4. Zmontowany zasilacz

rozmiarach, wyposażonych w przekładnie metalowe (<https://bit.ly/37pEFNb>). To bardzo dobre elementy. Są co prawda droższe i może nieco głośniejsze od innych, ale działają najpłynniej ze wszystkich i stanowią mój preferowany wybór, jeśli chodzi o serwo mechanizmy 9 g.

Szukałem również alternatywnych serwo mechanizmów 55 g i wybrałem typy Tower Pro MG996R (<https://bit.ly/3yC8Lsf>). Są to również serwo mechanizmy z przekładniami metalowymi. Wydaje się, że działają płynnie, a przy tym są niezbyt głośne.

Każda klasa serwo mechanizmów, które ocenilem, wydaje się mieć spójne wymiary i dlatego są one wymienne w stosowaniu. Nie mogę wypowiadać się na temat innych marek i modeli serwo mechanizmów o innych wymiarach.

Muszę też zauważyć, że znalezienie kart katalogowych serwo mechanizmów zawierających pełne informacje o opisywanych produktach okazało się sporym wyzwaniem”.

Masz dobry pomysł?

Nie wiem dlaczego, ale z jakiegoś powodu wyobrażam sobie, że czytasz ten artykuł i myślisz sobie: „Mam świetny pomysł dotyczący animatronicznej głowy Steve’a i Maxa, ale nie wiem, czy dam radę wysłać do Maxa e-mail i mu o tym powiedzieć”.

No cóż, może uda mi się skłonić Cię do wstania z wygodnego fotela, powolnego udania się do komputera i podzielenia się z nami swoim pomysłem. Już wiem – zorganizujemy konkurs! Szczegóły znajdziesz w następnym rozdziale czyli...

Zasilacz powraca

Obawiam się, że muszę rozpocząć tę część naszej dyskusji od stwierdzenia, że czuję się jak stary głupiec. „Ale co stary głupiec robiłby tutaj o tej porze dnia?” – słyszę wasze wołanie. Nie ma to jak stary dowcip.

Dlaczego jestem pełen rozgoryczenia i pogrążam się w zażenowaniu? No więc w poprzednim, 28. odcinku „Practical Electronics” z czerwca 2022 r. (EdW 1/2026) rozprawiałem o układzie zasilającym. Mój kumpel Joe Farr umieścił taki układ na płytce do programowania i testowania mikrokontrolerów PIC, którą zaprojektował dla niżej podpisanego. W ramach tej pracy Joe przygotował użyteczny układ samodzielnego zasilacza, który opisałem z całym bogactwem szczegółów, publikując między innymi projekt płytki drukowanej (rysunek 3a).

Być może pamiętacie, że Joe zaprojektował płytkę jako jednostronną, by ułatwić życie tym, którzy lubią wytrawiać płytki w domu. We wspomnianym artykule podałem też,

że Joe uprzejmie udostępnia do pobrania pliki projektowe.

Problem polega na tym, że kiedy nieco później rzuciłem okiem na projekt płytki, nagle zdałem sobie sprawę, że nie widzę żadnego miedzianego połączenia między elementami węzła 0 V (masy). Spójrzcie na przykład na dolne piny kondensatorów i sygnały 0 V na złączu SK2 (rysunek 3a).

Natychmiast wysłałem do Joe e-mail z pytaniem, dlaczego tak to zrobił i czy mamy wykonać te połączenia ręcznie. Kiedy już Joe podniósł się z podłogi i przestał się śmiać, wyjaśnił, że absolutnie nie ma potrzeby ręcznego wykonywania tych połączeń. Po prostu pominął rysunek tej części miedzi dla przejrzystości. Przesłał również inną wersję rysunku pokazującą całą miedź (rysunek 3b). To było tak oczywiste, że na pewno zdałbym sobie z tego sprawę, gdybym się chwilę zastanowił. I to wyjaśnia, dlaczego teraz czuję się tak głupio.

Wciągająca rywalizacja

Jak już wspomniałem, samodzielna płytka zasilacza to była rzecz, którą zrobiliśmy dla odprężenia. W ramach realizacji projektu Joe zlecił wykonanie pięciu płytek, z których jedną zmontował dla mnie, abym mógł się nią pobawić (rysunek 4). A to oznacza, że pozostały nam cztery nieobsadzone płytki.

Joe zasugerował, abyśmy zaofiarowali te płytki jako nagrody w małym konkursie (dotyczącym projektu głów robota; przypis redaktora). Omówiłem to z wybitnym wydawcą „Practical Electronics”, Mattem Pulzerem, chodzącą legendą w swojej dziedzinie. Uzgodniliśmy, że powinny być cztery kategorie (po jednej na płytkę), jak następuje:

- Najlepszy dodatek mechaniczny
- Najlepszy pod względem wizualnym efekt świetlny
- Najciekawszy przykład zastosowania czujników

I co dalej?

Pierwszy punkt, jak sądzę, mówi sam za siebie. Co do najciekawszego efektu świetlnego – należy pamiętać, że w każdym oku mamy 45 indywidualnie sterowanych trójkolorowych diod LED (rysunek 5), więc wyobrażacie sobie jaki niesamowity efekt można było wymyślić (wystarczyło podać jego opis, ponieważ realizacja to już było moje zadanie). Jeśli chodzi o najciekawszy sposób wykorzystania czujników, to myślimy o zastosowaniu jakiegoś czujnika do wykrywania tego, co dzieje się w otoczeniu głowy, i o tym, aby głowa w jakiś sposób na to reagowała (i znów wystarczyło to tylko opisać, nie trzeba było realizować). Ostatnia kategoria była



Rysunek 5. Jedna z gałek ocznych animatronicznej głowy. Zdjęcie: Steve Manley

całkowicie otwarta – co powinniśmy zrobić w tym projekcie w przyszłości?

Zadaniem czytelników, o ile zdecydowali się wziąć udział w konkursie, było przesłanie mi e-mailem swoich pomysłów. W każdej kategorii Autor pomysłu najbardziej przypadającego mi do gustu miał otrzymać jedną z nieobsadzonych płytek zasilacza.

Świat jest mały

Kiedy byłem dzieckiem, często słyszałem, jak starsi mówili: „świat jest mały”. Mówi się tak, by wyrazić zaskoczenie, gdy w nieoczekiwanym miejscu spotykamy kogoś znajomego. Inny przypadek to gdy podczas rozmowy z kimś odkrywasz, że macie wspólnego przyjaciela czy kolegę. No cóż, im jestem starszy, tym bardziej powiedzenie to się sprawdza.

Na przykład – w ramach projektu, nad którym obecnie pracuję, otrzymałem niedawno grawerkę/wycinarkę laserową OMTECH 40 W CO2 model DF0812-40BG (<https://amzn.to/3vQqMRG>). To naprawdę całkiem niezłe urządzenie. Bardzo chciałbym skorzystać z okazji i przekazać Wam, że instrukcja obsługi dołączona do tej maszyny była niezwykle pomocząca i pomogła mi szybko i łatwo rozpocząć pracę. Niestety nie mogę tego powiedzieć, ponieważ w rzeczywistości... nie dołączono żadnej instrukcji obsługi do tej maszyny. Szczerze mówiąc, dopiero później instrukcję znalazłem na stronie internetowej producenta.

Mam przyjaciela (prześciana się śmiać, naprawdę), którego nazwijmy Rick... no bo tak ma na imię. Rick mieszka w Birmingham w stanie Alabama, około 150 kilometrów na południe od miasta, które określam jako dom. Rick niedawno przeszedł na emeryturę, a ponieważ jest świetny w sprawach technicznych, zapytałem go, czy nie chciałby wybrać się w podróż, aby pomóc mi wszystko skonfigurować.

Całość zajęła nam tylko kilka godzin i wymagała paru wizyt w lokalnym markecie technicznym, gdzie kupiliśmy trochę

drobiazgów – takich jak 20-litrowe wiadro na wodę destylowaną do chłodzenia lasera (jedna z rzeczy, których się nie spodziewałem). Sprzętowa część operacji była gotowa do działania.

Z krótkich poszukiwań w Internecie wynikało, że oprogramowanie dostarczone z wycinarką nie spełni naszych oczekiwań. Na szczęście był wtorek, godzina 18., a więc lokalna „przestrzeń twórcza” (miejsce spotkań majsterkowiczów; przypis redaktora) – Makers Local 256 – była otwarta dla publiczności (<https://bit.ly/3sfTPMc>). Postanowiliśmy tam pójść i sprawdzić, czy jest jakiś ekspert od laserów, który mógłby nam doradzić. Ekspert był! W przyszłym tygodniu odwiedzi mnie w biurze i pomoże w kwestii oprogramowania. Więcej na ten temat opowiem w następnym artykule.

Powód, dla którego o tym tutaj piszę, jest taki, że w drodze powrotnej z „przestrzeni twórczej” opowiedziałem Rickowi o moich ostatnich historiach z PIC-ami. Jak wspomniałem w poprzednim artykule, Joe z Wielkiej Brytanii zapoznał mnie ze zintegrowanym środowiskiem programistycznym Proton IDE i kompilatorami Positron PIC BASIC, stworzonymi przez mojego nowego bohatera, Lesa Johnsona. Mówię „kompilatory” w liczbie mnogiej, bo Les dostarcza kompilatory oddzielne dla układów 8- i 16-bitowych, ale oba są zawarte w jednej cenie około 160 złotych. IDE automatycznie używa odpowiedniego kompilatora w zależności od wybranego układu PIC.

Chciałem powiedzieć, że minęło sporo czasu, odkąd programowałem w języku BASIC, ale właśnie przypomniałem sobie, że – jak wspominałem na jednym z moich blogów artykuł (<https://bit.ly/3KTWtHJO>) – kilka lat temu kupiłem nowoczesną wersję legendarnego Commodore 64 o nazwie THEC64 (<https://bit.ly/3smaoGu>). Replika ta oferuje wyjście wideo HDMI o wysokiej rozdzielczości, joystick w klasycznym stylu i 64 wbudowane klasyczne gry, a ponadto umożliwia programowanie w języku BASIC. Ale odbiegamy od tematu...

Chodzi o to, że kiedy wspominałem Rickowi, że odkryłem niesamowitą moc kompilatorów Positron PIC BASIC, od razu odpowiedział, że sam jest ich wielbiciele, że dobrze zna Lesa Johnsona i że regularnie komunikują się za pośrednictwem forum Positron BASIC Community (<https://bit.ly/3vRjRHO>). Cóż mogę powiedzieć poza tym, że świat naprawdę jest mały?

OK, teraz Wasza kolej. Pamiętajcie, że bardzo chętnie otrzymuję od Was komentarze, pytania i sugestie. Do następnego razu, życzę miłego dnia! ■

Clive „Max” Maxfield

AT-AD269S

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



AT-AD409PRO

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



BESTSELLERY sklepu AVT – sklep.avt.pl

Mikroskopy cyfrowe dla elektroników

Rabat dla Czytelników EdW
przy zakupie podaj kod **EdW2505MC**

Kod ważny do 30.09.2025

-3%

Rabat dla Prenumeratorów EdW
przy zakupie podaj numer prenumeraty

-6%

AT-AD246S-M

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



AT-AD407

Mikroskop cyfrowy 7 cali,
powiększenie do 270×
ANDONSTAR AD407



AT-AD249S-M

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



AT-AD210

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



Patronat EdW nad szkołami i uczelnianymi Kołami Naukowymi rozkwita i daje redakcji EdW impulsy zachęcające do wspierania edukacji szkolnej i uczelnianej. Działa sprzężenie zwrotne. Dostajemy mnóstwo wiadomości od uczniów, nauczycieli i studentów. Dla nich jest ta rubryka.



Kwantyzacja polega na przypisaniu chwilowej wartości sygnału analogowego do jednego z dyskretnych poziomów amplitudy, reprezentowanych przez kody cyfrowe. Powoduje to zaokrąglenie wartości sygnału i wprowadza nieunikniony błąd kwantyzacji.

Cyfrowe kontra analogowe

Cyfrowe przetwarzanie i obróbka analogowych sygnałów elektronicznych oferują wiele istotnych zalet. Przede wszystkim sygnał, od momentu gdy występuje w postaci cyfrowej, może być przesyłany i kopiowany w zasadzie bez ograniczeń, bez jakiegokolwiek utraty jakości. Kody cyfrowe nie są bowiem podatne na szумы i zniekształcenia, z których słyną analogowe systemy transmisji i odtwarzania.

Po drugie, na kodach cyfrowych można wykonywać złożone algorytmy matematyczne, co umożliwia przetwarzanie danych na najrozszybsze, często bardzo nietypowe sposoby. Przykładowo, w sposób całkowicie cyfrowy można podzielić pełne pasmo audio od 20 Hz do 20 kHz na bardzo strome i wąskie podzakresy częstotliwości, a następnie w każdym z nich realizować różnego rodzaju operacje obróbki sygnału. Później, ponownie w domenie cyfrowej, wąskie podzakresy można złożyć z powrotem w jedno szerokie pasmo audio. Tego rodzaju technik nie da się zrealizować za pomocą tradycyjnych układów analogowych.

Transformacje, które w elektronice analogowej są niewykonalne lub możliwe do wykonania jedynie z dużym nakładem pracy, w technice cyfrowej mogą być realizowane szybko i z dużą dokładnością. Złożone obliczenia są przy tym wykonywane przez procesory z bardzo dużą prędkością.

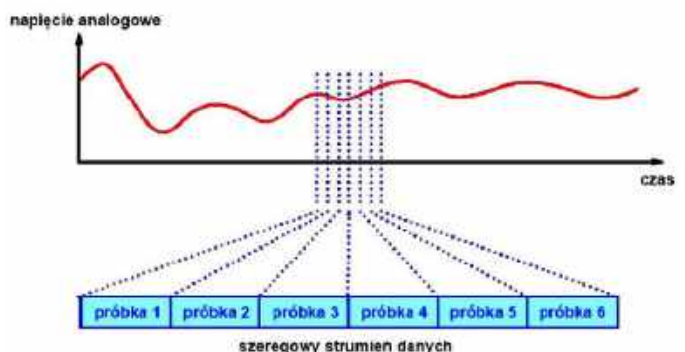
Zjawiska fizyczne są z natury analogowe

Większość sygnałów elektronicznych stanowi reprezentację zjawisk fizycznych i z tego powodu nie są to sygnały cyfrowe, lecz analogowe. Jako przykład można podać dźwięk. Powstaje on wówczas, gdy źródło – na przykład instrument muzyczny – powoduje sprężanie cząsteczek powietrza, a następnie ich rozprężanie. Powstała fala ciśnienia rozchodzi się w powietrzu i w innym miejscu wprawia w drgania błonę bębniową słuchacza.

Za pomocą mikrofonu zjawiska te można bardzo dokładnie przekształcić w elektryczne napięcia zmienne. Napięcia te mają charakter analogowy, co oznacza, że ich chwilowa wartość może przyjmować dowolną wartość pomiędzy pewnym minimum a maksimum. Na wykresie przedstawionym na poniższym rysunku, w górnej części, pokazano sygnał analogowy (czerwona krzywa). Z wykresu tego, który przedstawia przebieg napięcia w funkcji czasu, jednoznacznie wynika, że sygnał ten może przyjmować dowolną wartość pomiędzy dwiema granicami.

Od analogu do cyfry poprzez próbkowanie (sampling)

Sygnał analogowy jest przekształcany w kody cyfrowe. Odbывается to za pomocą przetwornika analogowo-cyfrowego (ADC). Taki układ w określonych chwilach czasu



Rysunek 1. Próbkowanie sygnału analogowego – chwilowe wartości napięcia są zamieniane na próbki w dyskretnych momentach czasu © 2020 Jos Verstraten

pobiera z sygnału analogowego „próbki” (ang. samples). Chwilowa wartość analogowa każdej z tych próbek zostaje następnie zamieniona na kod cyfrowy, którego postać stanowi miarę wielkości sygnału analogowego w danym momencie.

Próbkowanie czy kwantyzacja

Proces ten nazywa się „próbkowaniem sygnału analogowego” albo „kwantyzacją”. Jest on sterowany sygnałem zegarowym, który określa, ile próbek pobieranych jest w ciągu sekundy.

Przypis redaktora: Proces przekształcania sygnału analogowego do postaci cyfrowej obejmuje dwa odrębne etapy: próbkowanie (pobieranie próbek w dyskretnych chwilach czasu) oraz kwantyzację (przypisanie każdej próbce jednej z dyskretnych wartości amplitudy, a następnie jej zakodowanie). Częstotliwość zegara określa, jak często wykonywane jest próbkowanie (ile próbek na sekundę). Autor najpewniej stosuje skrót myślowy, wedle którego próbkowanie jest oczywistym krokiem poprzedzającym kwantyzację.

Kwantyzacja jest techniką o fundamentalnym znaczeniu: fotografie, muzyka oraz wideo są obecnie niemal wyłącznie przesyłane i przetwarzane w postaci cyfrowej.

Od cyfry do analogu

Jeżeli zapisane próbki cyfrowe mają zostać ponownie przekształcone w sygnał analogowy, konieczne jest zastosowanie przetwornika cyfrowo-analogowego (DAC). Układ ten zamienia kody cyfrowe z powrotem na sygnał analogowy.

System podstawowy

Podstawowy system cyfrowego przetwarzania sygnałów analogowych przedstawiono na poniższym rysunku. W zaprezentowanym przykładzie system ten służy do opóźniania sygnału analogowego o określony czas.

Sygnał wejściowy jest za pomocą przetwornika ADC zamieniany na zestaw kodów cyfrowych od Qa do Qd. Chwilowa postać tych kodów stanowi miarę chwilowej wartości sygnału analogowego. Należy przy tym zaznaczyć, że pokazanej postaci kodów nie należy przypisywać żadnego szczególnego znaczenia – jest to jedynie przykład mający na celu zilustrowanie zasady działania. Przekształcanie sygnału analogowego w kody cyfrowe nie odbywa się w sposób losowy, lecz w regularnych odstępach czasu Δt . Sygnały cyfrowe (w tym prostym przykładzie cztery) są następnie poddawane cyfrowemu przetwarzaniu. W omawianym przypadku cztery sygnały cyfrowe trafiają do cyfrowego układu opóźniającego.

Gdy dane są ponownie potrzebne, są one odczytywane z systemu cyfrowego w rytmie tego samego zegara. Sygnały wyjściowe Qe do Qh zostają następnie przekształcone w przetworniku DAC z powrotem w sygnał analogowy. Sygnał ten nazywany jest odzyskanym sygnałem analogowym.

Przybliżenie kwantyzacyjne

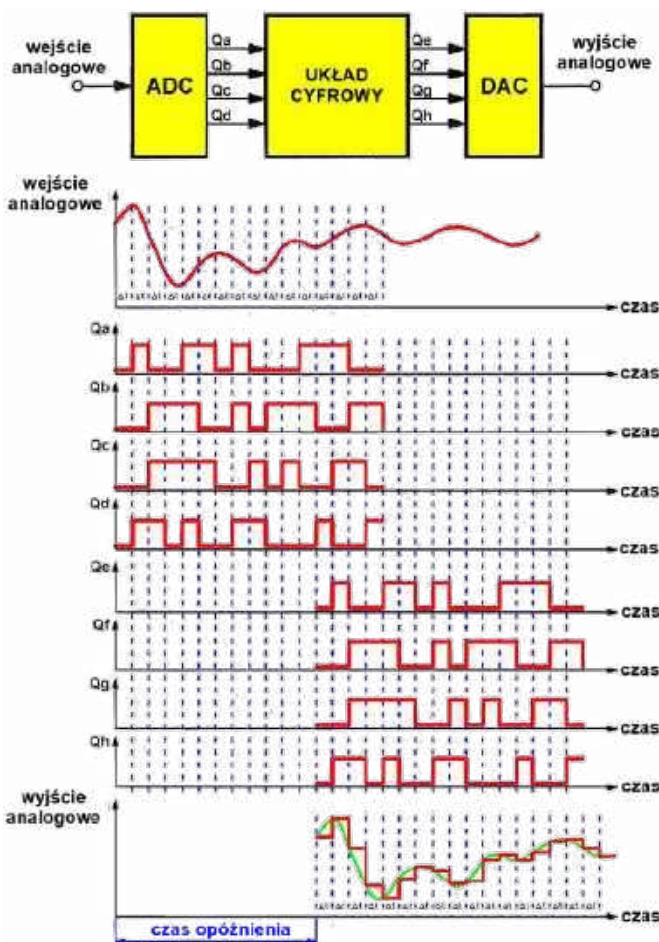
Sygnał cyfrowy jest z definicji sygnałem, który może przyjmować tylko dwie wartości: „L” lub „H”. Jest więc oczywiste, że również kombinacja kilku sygnałów cyfrowych może zawierać jedynie ograniczoną liczbę możliwych kombinacji stanów „L” i „H”.

Jeżeli sygnał analogowy zostanie przekształcony w kod cyfrowy – niezależnie od tego, jak bardzo złożona jest jego struktura – kod ten będzie mógł jedynie przybliżać chwilową wartość sygnału analogowego. **Nieskończona liczba możliwych wartości sygnału analogowego zostaje w ten sposób odwzorowana na skończoną liczbę kombinacji kodów cyfrowych.**

Różnica ta pomiędzy przetwarzaniem sygnałów analogowych i cyfrowych stanowi jedną z najważniejszych cech systemów ADC i DAC. Zjawisko to określa się mianem przybliżenia kwantyzacyjnego. Kody wyjściowe przetwornika ADC są w każdej chwili jedynie cyfrowym przybliżeniem analogowego sygnału wejściowego.

Nietrudno zauważyć, że jeżeli te przybliżone kody zostaną następnie, za pomocą przetwornika DAC, przekształcone z powrotem w sygnał analogowy, to odzyskany sygnał również będzie jedynie przybliżeniem sygnału pierwotnego. Z tego powodu mówi się, że kwantyzacja zawsze prowadzi jedynie do schodkowego przybliżenia oryginalnego sygnału analogowego.

Na przedstawionym wcześniej rysunku schodkowy charakter tego przybliżenia jest dobrze widoczny, a na kolejnym



Rysunek 2. Podstawowy system cyfrowego przetwarzania sygnału analogowego © 2020 Jos Verstraten

rysunku został on – dla większej czytelności – zaprezentowany w sposób celowo nieco przesadzony. W praktycznych zastosowaniach pracuje się oczywiście z więcej niż czterema sygnałami cyfrowymi oraz z dużo większą liczbą próbek pobieranych w ciągu sekundy.

Zniekształcenia kwantyzacji, czyli szum kwantyzacji

Pomiędzy oryginalnym sygnałem analogowym a odzyskanym sygnałem analogowym zawsze będzie występować pewna różnica, czyli zniekształcenie. Zniekształcenie to można wprawdzie ograniczać, jednak nigdy nie da się go całkowicie wyeliminować. Sygnały analogowe, które po przejściu przez systemy ADC + DAC są ponownie zamieniane na sygnały analogowe, z definicji ulegają zniekształceniu.

To zniekształcenie, wynikające bezpośrednio z zasady działania układów kwantyzujących, nazywane jest zniekształceniem kwantyzacji. Ponieważ w praktyce objawia się ono najczęściej w postaci drobnych sygnałów o częstotliwości podstawowej równej częstotliwości zegara, bywa ono również określane mianem szumu kwantyzacji. Zegar pracuje bowiem z wysoką częstotliwością, a sygnały o wysokich częstotliwościach są często utożsamiane z szumem.

Przypis redaktora: W ujęciu teoretycznym szum kwantyzacji ma charakter szerokopasmowy. Składowe powiązane z częstotliwością zegara mogą pojawiać się w praktycznych przetwornikach jako efekt ich niedoskonałości.

Szum kwantyzacji stanowi jeden z najważniejszych parametrów charakteryzujących systemy cyfrowego przetwarzania sygnałów, a jego minimalizacja jest jednym z największych wyzwań stojących przed każdym elektronikiem zajmującym się projektowaniem układów kwantyzujących.

Przykład szumu kwantyzacji

Jeżeli dysponujesz oscyloskopem cyfrowym, szum kwantyzacji można zauważyć na niemal każdym obrazie wyświetlanym na ekranie.

Na poniższej fotografii pokazano, w jaki sposób takie urządzenie przedstawia czyste napięcie sinusoidalne w sytuacji, gdy zarówno wzmocnienie, jak i podstawa czasu zostaną znacznie zwiększone.

Warto dodać, że niemal wszystkie nowoczesne oscyloskopy cyfrowe wyposażone są w pewien trik pozwalający ograniczyć ten szum – funkcję Average Sampling.

Parametry systemów kwantyzacji

Długość słowa

Długość słowa określa liczbę bitów wykorzystywanych do przekształcania sygnału analogowego w kod cyfrowy. Jeżeli pracujesz z systemem szesnastobitowym (co jest często spotykanym standardem), mówi się o długości słowa równej 16 bitów. Bez dodatkowych wyjaśnień oczywiste jest, że wraz ze wzrostem liczby bitów maleje poziom szumu kwantyzacji. Z tego powodu jakość obrazu w dziesięciobitowym oscyloskopie cyfrowym jest znacznie lepsza niż w jego ośmiobitowym odpowiedniku.

Rozdzielczość

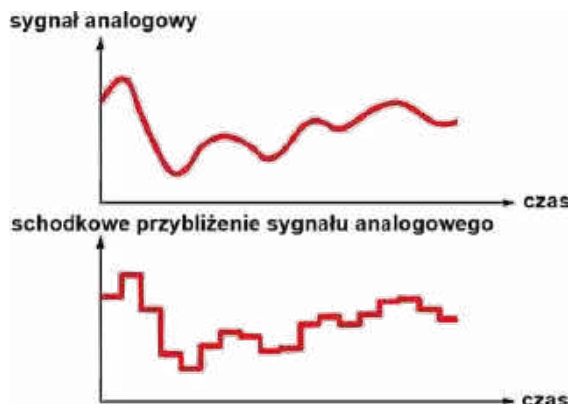
Rozdzielczość przetwornika ADC określa, na ile „stref kwantyzacji” można podzielić sygnał analogowy po jego digitalizacji. Oczywiście jest, że system o długości słowa równej 1 bitowi ma rozdzielczość 2. Rozdzielczość można obliczyć, podnosząc liczbę 2 do potęgi równej liczbie bitów. System o długości słowa 16 bitów ma zatem rozdzielczość $2^{16}=65\,536$. W takim systemie chwilową wartość sygnału analogowego można przyporządkować do 65 536 stref kwantyzacji.

Zakres dynamiczny

Zakres dynamiczny przedstawia logarytmiczną relację pomiędzy minimalną a maksymalną zmianą sygnału, jaką przetwornik DAC może wywołać w odzyskanym sygnale analogowym. Maksymalna zmiana występuje wówczas, gdy wszystkie bity jednocześnie przełączają się ze stanu „L” na „H”. Minimalna zmiana pojawia się wtedy, gdy ze stanu „L” na „H” przełącza się jedynie bit LSB.

Oczywiste jest, że zakres dynamiczny rośnie wraz ze wzrostem liczby bitów biorących udział w procesie konwersji. Im większa długość słowa, tym większy zakres dynamiczny. Zakres dynamiczny wyrażany jest w decybelach (dB) i bywa również określany mianem stosunku sygnału do szumu systemu.

Przypis redaktora: W idealnym przetworniku, w którym dominującym źródłem szumu jest szum kwantyzacji, zakres dynamiczny



Rysunek 3. Schodkowe przybliżenie jest podstawową cechą kwantyzacji © 2020 Jos Verstraten



Rysunek 4. Szum kwantyzacji widoczny na ekranie oscyloskopu cyfrowego © 2020 Jos Verstraten

i stosunek sygnału do szumu (SNR) są ściśle powiązane z liczbą bitów. W praktycznych układach pojęcia te nie zawsze są tożsame, ponieważ na SNR wpływają również inne źródła szumu i zniekształceń.

Wielkość bitu

Wielkość bitu określa, o ile woltów zmienia się odzyskany sygnał analogowy na wyjściu przetwornika DAC w chwili, gdy bit LSB przełącza się ze stanu „L” na „H”. Wielkość bitu oraz zakres dynamiczny są więc w istocie dwiema wielkościami opisującymi to samo zjawisko fizyczne.

O ile jednak zakres dynamiczny opisuje pewną relację i jest niezależny od maksymalnej wartości napięcia, które może zostać odtworzone, o tyle podawanie wielkości bitu ma sens jedynie wtedy, gdy znana jest maksymalna wartość napięcia wyjściowego przetwornika DAC. Często wartość bitu w danym systemie określa się po prostu jako LSB systemu.

Zależność między długością słowa, rozdzielczością, zakresem dynamicznym i wartością bitu

Logiczne jest, że pomiędzy czterema omówionymi podstawowymi wielkościami każdego systemu kwantyzacji istnieje związek przyczynowy. Zależność tę przedstawiono w tabeli zamieszczonej na poniższym rysunku dla długości słowa od 1 do 16 bitów. Podane wartości wielkości bitu odnoszą się do sygnału analogowego o maksymalnym napięciu równym 10 V.

Długość słowa	Rozdzielczość	Zakres dynamiczny	Wielkość bitu przy 10 V max.
1	2	6 dB	5,0 V
2	4	12 dB	2,5 V
4	16	24 dB	625 mV
8	256	48 dB	39,1 mV
12	4096	72 dB	2,4 mV
16	65536	96 dB	0,152 mV

Rysunek 5. Zestawienie najważniejszych parametrów © 2020 Jos Verstraten

Częstotliwość próbkowania

Digitalizacja sygnału analogowego nie jest procesem ciągłym, lecz procesem sterowanym zewnętrznym impulsem zegarowym. Częstotliwość tego sygnału zegarowego nazywana jest częstotliwością próbkowania. Określa ona, ile próbek napięcia analogowego na wejściu jest pobieranych w ciągu sekundy oraz ile różnych kodów cyfrowych pojawi się w ciągu sekundy na wyjściu przetwornika ADC. Jednostką tej wielkości są próbki na sekundę, oznaczane jako Sa/s. Ponieważ obecnie często stosuje się częstotliwości próbkowania rzędu megaherców, jednostka ta bywa zapisywana jako MSa/s, a nawet GSa/s, czyli odpowiednio milion lub miliard próbek na sekundę.

Waga bitów oraz MSB i LSB

W poniższym przykładzie przedstawiono analogowe napięcie rosnące w postaci piły, o maksymalnej wartości 15 V. Napięcie to jest kwantyzowane za pomocą przetwornika ADC o rozdzielczości czterech bitów. Układ ten generuje standardowy kod binarny. Częstotliwość zegara jest równa szesnastokrotności częstotliwości przebiegu piłokształtnego. W tabeli pokazano stan czterech bitów przy każdym impulsie zegara.

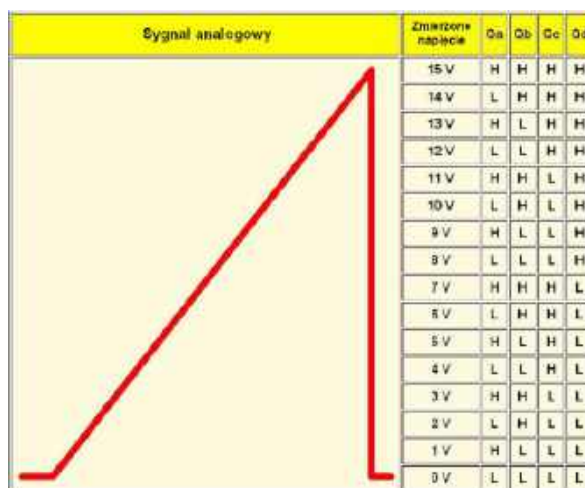
Zauważalne jest, że bit Qa zmienia swój stan przy każdym wzroście napięcia piły o 1 V. Bit Qb zmienia stan po każdym wzroście napięcia o 2 V. Cztery bit, Qd, zmienia swój stan tylko jeden raz – przy napięciu piły równym 8 V. Oznacza to, że każdemu bitowi można przypisać określoną „wagę”. Bit Qa jest bitem „najlżejszym”, natomiast bit Qd – „najcięższym”. Zmiana stanu bitu Qa odpowiada bowiem zmianie napięcia analogowego o 1 V, podczas gdy zmiana stanu bitu Qd odpowiada zmianie napięcia analogowego o 8 V.

Z tego powodu bit Qa nazywany jest Least Significant Bit, w skrócie LSB, natomiast bit Qd określany jest jako Most Significant Bit, w skrócie MSB.

Twierdzenie o próbkowaniu Minimalna częstotliwość próbkowania ma kluczowe znaczenie

Istotnym pytaniem jest, ile próbek należy pobierać w ciągu sekundy. Oczywiście jest, że jakość systemu poprawia się wraz ze wzrostem liczby próbek pobieranych na sekundę. Większa liczba próbek oznacza jednak wyższe częstotliwości zegara, a tym samym bardziej złożoną i droższą elektronikę.

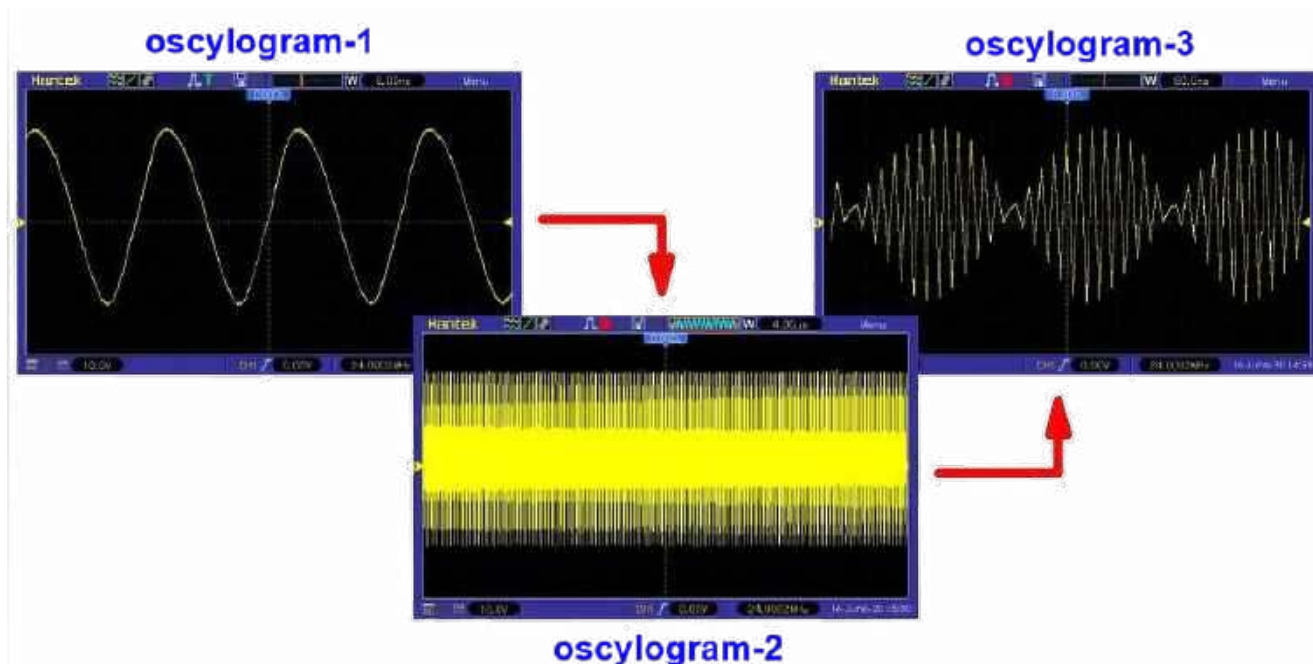
Z tego powodu interesujące jest przede wszystkim to, jak niewielką liczbę próbek można pobierać, aby mimo to mieć pewność, że odzyskany sygnał analogowy będzie przypominał sygnał wejściowy przetwornika ADC.



Rysunek 6. Wyjaśnienie pojęć LSB i MSB © 2020 Jos Verstraten

Przykład z praktyki

Z pewnej „czarnej skrzynki” wyprowadzony jest nieznaną sygnał. Pojawia się pytanie, jaki jest jego rzeczywisty przebieg. Jeżeli do wyjścia podłączysz oscyloskop cyfrowy i ustawisz szybką podstawę czasu, na przykład 8 ns/div, sygnał będzie wyglądał tak, jak pokazano na **oscylogramie 1** na poniższej fotografii. Czarna skrzynka zdaje się więc generować sinusoidę o częstotliwości 24 MHz. Jeżeli jednak zmniejszysz szybkość podstawy czasu, na przykład do 4 μ s/div, na ekranie pojawi się zupełnie inny obraz – patrz **oscylogram 2**. Sprawia on wrażenie, jakby w sygnale występowały różnego rodzaju nieregularności. Naturalnie chcesz przyjrzeć się temu zjawisku dokładniej, dlatego zatrzymujesz oscylogram 2 na ekranie i rozciągasz obraz do 80 ns/div. Wówczas na ekranie pojawia się **oscylogram 3**. Powstaje pytanie, jaki sygnał rzeczywiście generuje czarna skrzynka: ten z oscylogramu 1 czy ten z oscylogramu 3.



Rysunek 7. Twierdzenie o próbkowaniu w praktyce © 2020 Jos Verstraten

Na czym polega różnica? Dwa pierwsze obrazy zostały zarejestrowane przy tej samej liczbie próbek na sekundę. W drugim przypadku na ekranie wyświetlanych jest jednak znacznie więcej okresów sygnału niż w pierwszym obrazie. Oznacza to, że na jeden okres przypada znacznie mniej próbek niż w obrazie 1. Jeżeli teraz obraz 2 zostanie rozciągnięty tak, aby na ekranie pozostało jedynie kilkadziesiąt okresów, okaże się, że z powodu ograniczonej liczby próbek przypadających na jeden okres może powstać całkowicie błędny obraz rzeczywistego sygnału.

Twierdzenie o próbkowaniu

Jak widać, nie można w nieskończoność zmniejszać częstotliwości próbkowania. Z matematycznego punktu widzenia można wykazać, że częstotliwość próbkowania musi być co najmniej dwukrotnie większa od najwyższej częstotliwości występującej w sygnale analogowym. To ogólne i niezwykle istotne prawo znane jest jako twierdzenie o próbkowaniu. Słowo „twierdzenie” jest po prostu bardziej uczonym określeniem powszechnie znanego pojęcia „teza” lub „reguła”.

Jeżeli próbkowanie odbywa się z częstotliwością niższą, wówczas odtworzenie kształtu sygnału analogowego na podstawie kolejnych próbek cyfrowych staje się całkowicie niemożliwe. Przeprowadzony wcześniej eksperyment z użyciem oscyloskopu cyfrowego stanowił tego bardzo dobry przykład.

Graficzna ilustracja twierdzenia o próbkowaniu

Chociaż w ramach tego artykułu przedstawienie matematycznego dowodu byłoby zbyt obszerne, efekt ten można w bardzo wyrazisty sposób zilustrować graficznie. Na poniższym rysunku pokazano sinusoidalne napięcie analogowe o częstotliwości 15 kHz, próbkowane sygnałem zegarowym o częstotliwości stopniowo zmniejszanej. Jeżeli następnie kolejne kody cyfrowe zostaną przekształcone za pomocą przetwornika DAC w napięcie analogowe, powstaje sygnał, który powinien stanowić przybliżenie przebiegu sinusoidalnego.

Daje to poprawny rezultat przy próbkowaniu z częstotliwością co najmniej 30 kHz (dla sygnału 15 kHz), a w praktyce stosuje się jeszcze pewien zapas. Gdy częstotliwość próbkowania spadnie poniżej 30 kHz, pojawiają się zniekształcenia aliasowe i odzyskany przebieg może mieć zupełnie inną (zwykle niższą) częstotliwość niż sygnał wejściowy. Przy bardzo niskich częstotliwościach próbkowania wynik przestaje przypominać oryginalną sinusoidę.

Podobne zjawisko wystąpiło również podczas próbkowania sygnału pochodzącego z „czarnej skrzynki”.

Częstotliwość aliasowa i zniekształcenia aliasowe

Tę znacznie niższą częstotliwość nazywa się częstotliwością aliasową, a powstające w takim przypadku podczas kwantyzacji duże zniekształcenie określa się mianem zniekształcenia aliasowego. Ryzyko wystąpienia zniekształceń aliasowych pojawia się zawsze wtedy, gdy próbkowanie odbywa się z częstotliwością niższą niż dwukrotność najwyższej częstotliwości występującej w sygnale.

Jeżeli sygnał wejściowy ma znaną częstotliwość, zniekształceń aliasowych można uniknąć, dobierając częstotliwość próbkowania co najmniej dwukrotnie wyższą.

Obrazowy przykład z filmu

Bardzo znanym przykładem zniekształceń aliasowych jest zjawisko, w którym na filmie koła pojazdu wydają się obracać wstecz. Szprychy kół poruszają się bowiem w rzeczywistości analogowej zbyt szybko w stosunku do cyfrowego próbkowania realizowanego przez kamerę filmową (24 klatki na sekundę), co prowadzi do wystąpienia tego rodzaju błędu aliasowego.

Harry Nyquist i częstotliwość Nyquista

Harry Nyquist wykazał matematycznie, że tego rodzaju błędy nie mogą wystąpić, jeżeli częstotliwość próbkowania jest większa niż dwukrotność najwyższej składowej częstotliwości występującej w próbkowanym sygnale. Częstotliwość tę nazywa się częstotliwością Nyquista.

Przypis redaktora: W terminologii stosowanej obecnie częstotliwość Nyquista jest równa połowie częstotliwości próbkowania. Warunek Nyquista mówi natomiast o konieczności próbkowania z częstotliwością co najmniej dwukrotnie wyższą od najwyższej składowej częstotliwości sygnału.

Podczas próbkowania sygnału należy za pomocą filtru dolnoprzepustowego zapewnić, aby wszystkie częstotliwości wyższe niż częstotliwość Nyquista zostały odfiltrowane. Tylko wówczas nie występuje ryzyko powstania zniekształceń aliasowych.

Przykładowo, w przypadku płyty audio CD próbkowanie odbywa się z częstotliwością zegara 44,1 kHz. Częstotliwość Nyquista wynosi w tym przypadku 22,05 kHz. Aby zapobiec zniekształceniom aliasowym, filtr analogowy osłabia wszystkie częstotliwości powyżej 22,05 kHz jeszcze przed procesem kwantyzacji dźwięku, tak aby nie miały one już wpływu na wynik próbkowania.

W tym celu konieczne jest zastosowanie filtru dolnoprzepustowego o bardzo stromym zboczu tłumienia.

Schemat blokowy kwantyzacji sygnałów analogowych Więcej niż tylko ADC!

Na podstawie dotychczas przedstawionych informacji można określić, jakie elementy są niezbędne do przekształcenia sygnału analogowego w kod cyfrowy. Podstawowy schemat blokowy takiego systemu przedstawiono na poniższym rysunku.

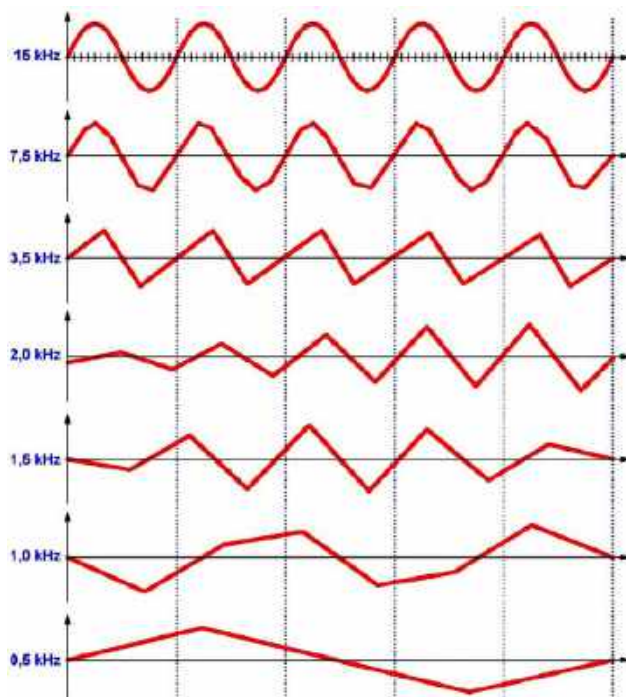
Sygnał przeznaczony do digitalizacji jest najpierw prowadzony przez filtr antyaliasingowy, a następnie trafia do układu określonego jako sample and hold (S&H). Wyjście tego układu jest połączone z analogowym wejściem przetwornika ADC. Zarówno układ S&H, jak i przetwornik ADC są sterowane zewnętrznym sygnałem zegarowym, który ustala częstotliwość próbkowania.

Filtr dolnoprzepustowy

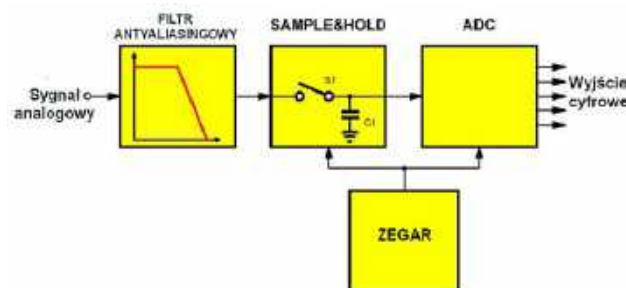
Filtr ten musi być filtrem o bardzo stromym zboczu tłumienia, który bez osłabienia przepuszcza maksymalną częstotliwość sygnału poddawanego kwantyzacji, natomiast częstotliwości dwukrotnie wyższe tłumi o kilkadziesiąt decybeli. W praktyce oznacza to konieczność zastosowania aktywnych filtrów trzeciego, a nawet piątego rzędu.

Konieczność zastosowania zewnętrznego zegara

Przetwornik ADC nie pracuje w sposób ciągły. Jeżeli pozwoli się, aby kod cyfrowy na jego wyjściu był na bieżąco dostosowywany do zmian analogowego sygnału wejściowego,



Rysunek 8. Graficzne objaśnienie twierdzenia o próbkowaniu © 2020 Jos Verstraten



Rysunek 9. Schemat blokowy układu kwantyzacji © 2020 Jos Verstraten

na wyjściu ADC może powstać prawdziwy chaos kodów. Proces przekształcania sygnału analogowego w kod cyfrowy wymaga pewnego czasu. Każdy przetwornik ADC charakteryzuje się określonym czasem konwersji, który zwykle mieści się w zakresie mikrosekund. Ponadto poszczególne bity nie reagują jednocześnie – jeden bit może zmieniać swój stan nieco szybciej niż inny.

Jeżeli bez zewnętrznej kontroli analogowe napięcie byłoby na bieżąco zamieniane na kody cyfrowe, system mógłby przypadkowo odczytać kod w chwili, gdy jeden lub kilka bitów dopiero dostosowuje się do nowej sytuacji na wejściu analogowym. Taki kod mógłby wówczas zawierać całkowicie błędną informację.

Aby uniknąć tych problemów, konieczne jest zastosowanie zewnętrznego zegara. Oznacza to, że w regularnych odstępach czasu pobierana jest próbka chwilowej wartości sygnału wejściowego, następnie próbka ta jest na krótko przechowywana w pamięci analogowej, po czym zostaje przekształcona w odpowiadający jej kod cyfrowy. Dzięki temu nie występują opisane zjawiska, ponieważ wartość próbki pozostaje stała. Po upływie czasu konwersji przetwornika ADC kod cyfrowy może zostać odczytany i dalej przetwarzany. Następnie można pobrać kolejną próbkę i powtórzyć cały proces.

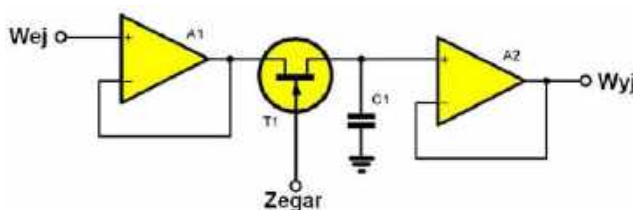
Układ sample and hold

Pamięć analogowa, w której próbka sygnału analogowego jest na krótko przechowywana, realizowana jest zawsze w postaci układu sample and hold, w skrócie S&H. Podstawową zasadę działania układu S&H przedstawiono na poniższym rysunku. Sygnał analogowy podawany jest na bufor A1. Wyjście tego bufora, poprzez przełącznik elektroniczny wykonany w postaci tranzystora FET T1, połączone jest z drugim buforem A2. Pomiędzy wejściem tego bufora a masą dołączony jest kondensator C1.

Po wysterowaniu przełącznika elektronicznego kondensator ładuje się do chwilowej wartości analogowego napięcia wejściowego. Ponieważ pojemność kondensatora jest bardzo mała, napięcie na nim może z bardzo małym opóźnieniem podążać za zmianami napięcia wejściowego. Gdy przełącznik elektroniczny zostanie otwarty, połączenie pomiędzy wyjściem pierwszego bufora a wejściem drugiego bufora zostaje przerwane. W efekcie napięcie na kondensatorze nie może się rozładować i pozostaje stałe.

Oczywiście pewna ilość ładunku będzie stopniowo upływać przez bardzo dużą impedancję wejściową drugiego bufora oraz przez różnego rodzaju rezystancje upływu występujące w układzie. Ponieważ jednak próbkowanie odbywa się stosunkowo szybko, można przyjąć, że w czasie trwania próbki napięcie na kondensatorze pozostaje równe chwilowej wartości sygnału wejściowego w momencie otwarcia przełącznika.

W tym momencie można zleczyć przetwornikowi ADC dokonanie kwantyzacji napięcia występującego na wyjściu układu S&H. Gdy kod cyfrowy na wyjściach przetwornika ADC stanie się stabilny, może on zostać odczytany, po czym można rozpocząć kolejny cykl próbkowania. ■



Rysunek 10. Podstawowy schemat układu sample and hold © 2020 Jos Verstraten

Jos Verstraten

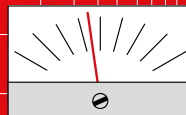
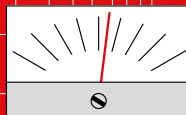
REKLAMA

m.technik

Ciekawy świata są zawsze młodzi

przejrzysz i kupisz w prezencie
na każdą okazję na stronie
www.ulubionykiosk.pl

AUDIO OUT

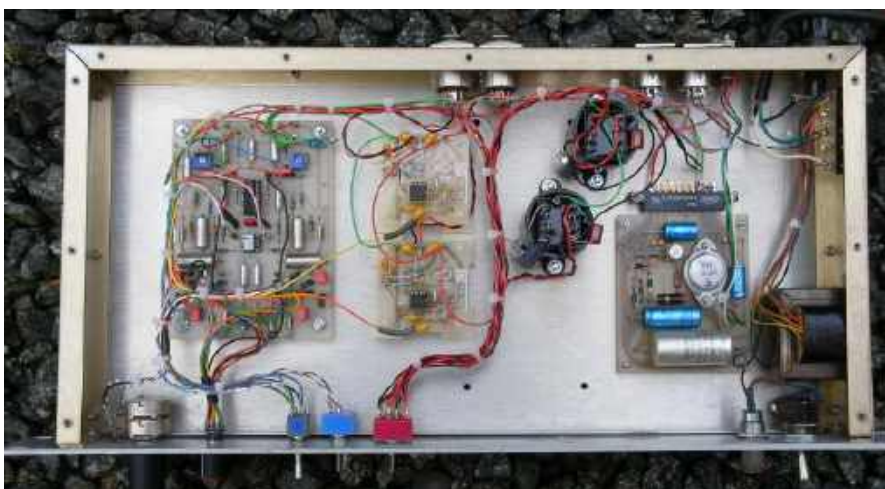


Uniwersalna płytka z pojedynczym wzmacniaczem operacyjnym, zoptymalizowana do zastosowań w układach akustycznych, część 1

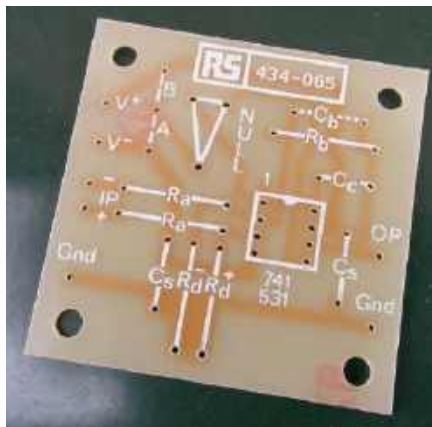
Często temat tej rubryki jest inicjowany przez Czytelników. W tym przypadku jeden z nich poszukiwał testera do wzmacniaczy operacyjnych. Stworzenie testera, który by mierzył wszystkie parametry wzmacniaczy operacyjnych byłoby trudne. Jest to zadanie dla dużego zespołu, nie dla mnie, dlatego ograniczę się do prostszych rozwiązań.

Zakres zadania

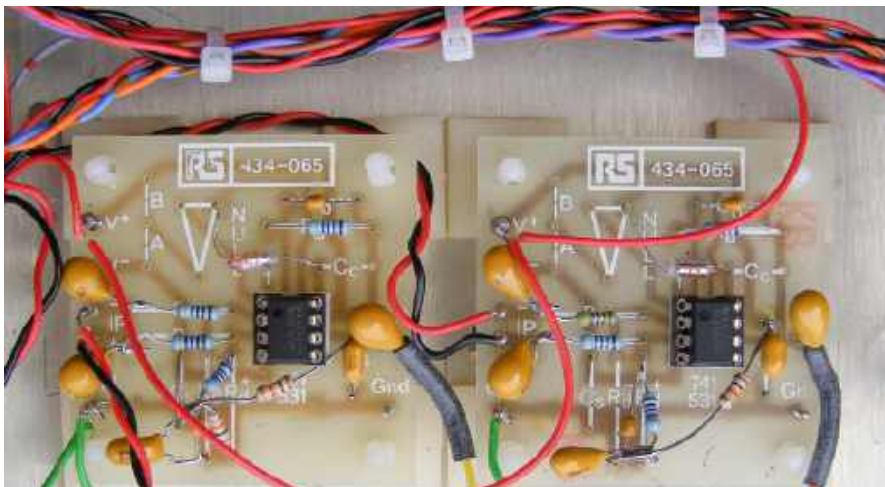
W sprzęcie akustycznym wzmacniacze operacyjne są zaskakująco niezawodne, a ich awaryjność jest prawie tak niska jak w przypadku pojedynczych tranzystorów małej mocy. Najczęstszą przyczyną awarii wzmacniacza operacyjnego jest uszkodzenie jego wejść – na przykład po podłączeniu do wejścia dużych, nierozładowanych kondensatorów lub przedostaniu się napięcia fantomowego o wartości 48 V w wyniku braku stosownych zabezpieczeń. Zazwyczaj skutkuje to trwałym uszkodzeniem struktury wewnętrznej układu. W takich przypadkach wystarczy przeprowadzić prosty test za pomocą multimetru lub sprawdzić, czy nie poszedł przysłowiowy dym. Trudniejszą do wykrycia



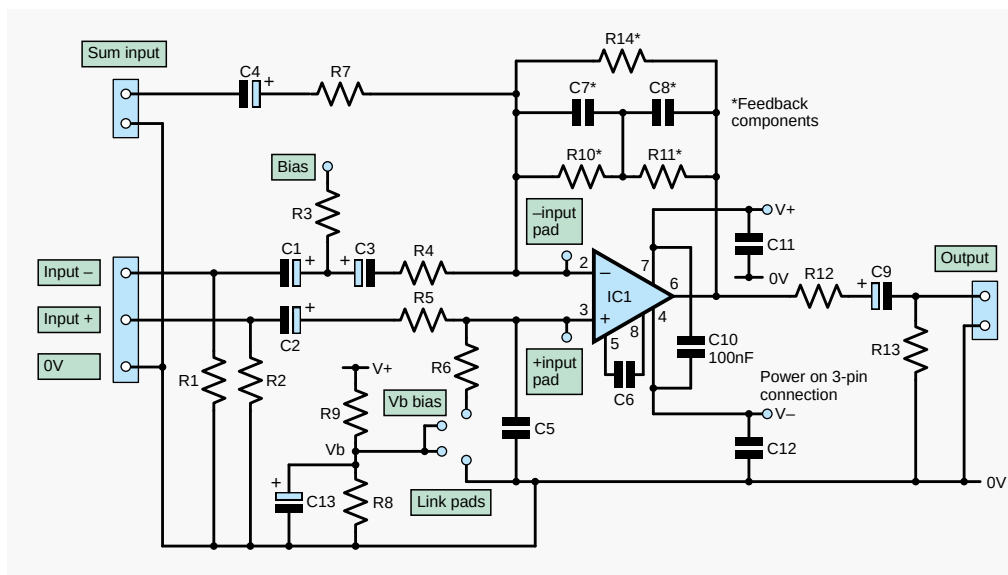
Rysunek 2. Typowe zastosowanie płytek RS w sprzęcie studyjnym. Do kompresora dynamiki uwypuklającego ciche dźwięki dodane zostały symetryczne wejścia XLR



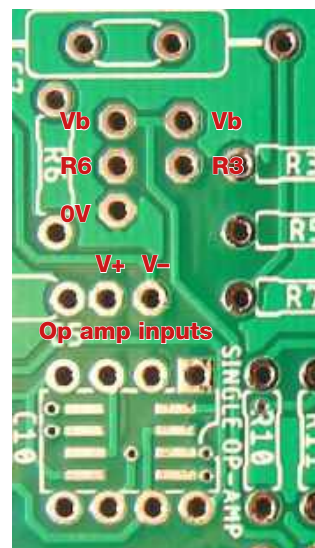
Rysunek 1. Wycofana z produkcji płytka drukowana firmy RS. Obecnie dostępny jest jej udoskonalony zamiennik, zoptymalizowany pod kątem pracy w układach akustycznych, jednak doskonale nadający się do innych zastosowań



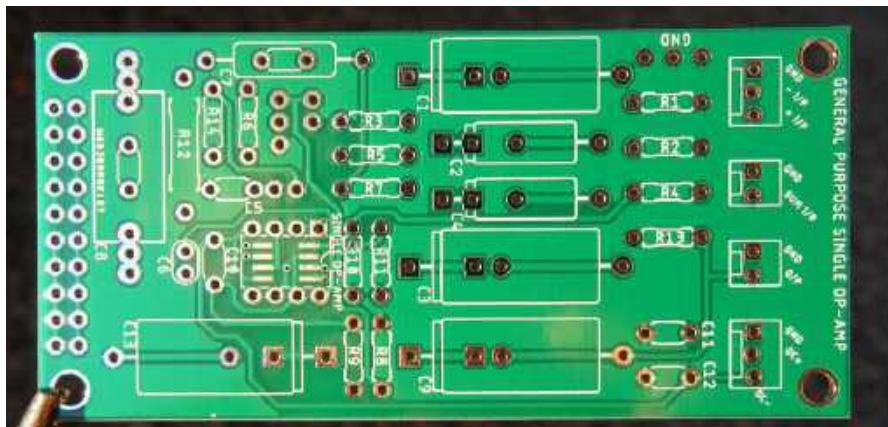
Rysunek 3. Problemem występującym na płytce RS był brak miejsca na duże kondensatory wymagane w układach akustycznych



Rysunek 4. Schemat kompletnego układu uniwersalnej płytki ze wzmacniaczem operacyjnym. Nie wszystkie elementy są używane jednocześnie, tylko kilka niezbędnych do realizacji konkretnego układu



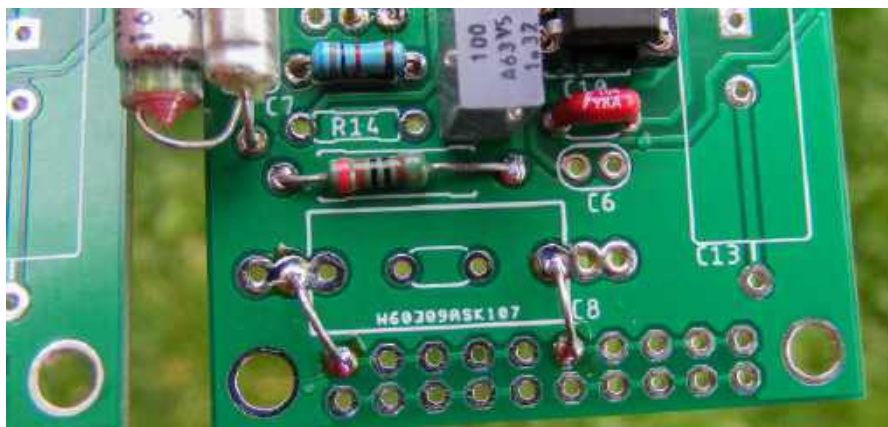
Rysunek 7. Na płytce dostępne są punkty lutownicze, pozwalające na polaryzację lub zerowanie wejść wzmacniacza operacyjnego. Istnieją również punkty prowadzące do pinów wejściowych 2 i 3 wzmacniacza, umieszczone w pobliżu kondensatora C5. Są one przydatne, gdy trzeba podłączyć potencjometry do regulacji wzmocnienia. Istnieje również punkt połączony z rezystorem R3, służącym do polaryzacji kondensatorów bipolarnych złożonych z elementów C1 i C3. Ponadto są dostępne trzy punkty zerujące – patrz rysunek 5



Rysunek 5. Widok płytki ze wzmacniaczem operacyjnym

awarią jest częściowe uszkodzenie tranzystorów wejściowych, powodujące wzrost poziomu szumów, co stanowi poważny problem w przypadku urządzeń akustycznych. Najlepszym sposobem sprawdzenia stanu

układu jest umieszczenie go we wzmacniaczu o wysokim wzmocnieniu i odsłuchanie sygnału wyjściowego, lub jego sprawdzenie za pomocą oscyloskopu. Ostatecznie zdecydowałem, że najlepszym sposobem

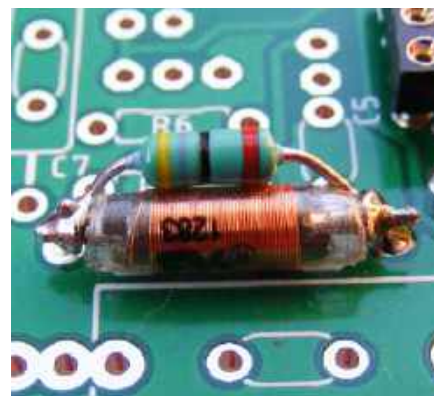


Rysunek 6. Na płytce umieszczone są cztery dodatkowe zestawy punktów lutowniczych, przeznaczone do eksperymentowania. Jest to pięć połączonych ze sobą punktów, podobnie jak w płytce prototypowej

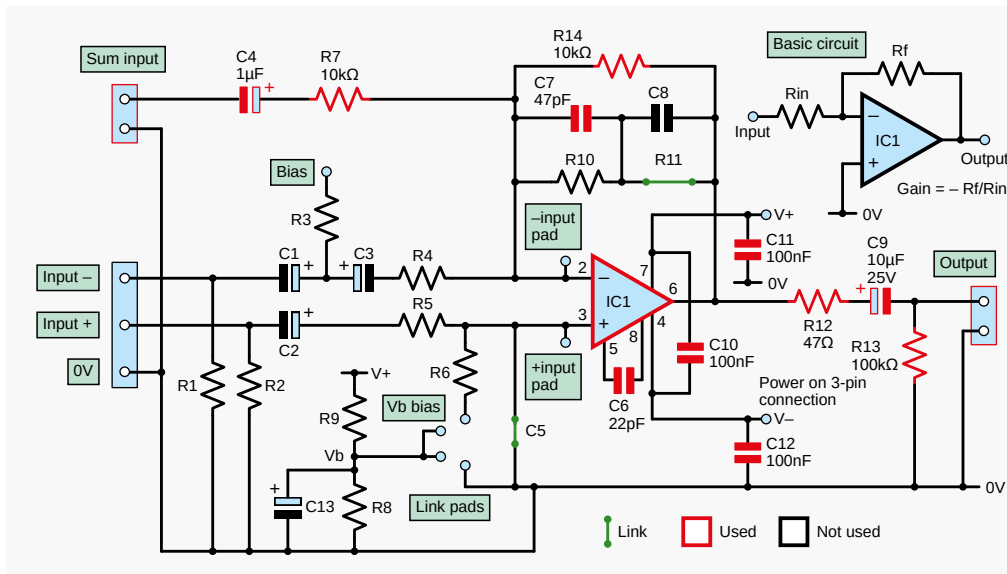
na testowanie wzmacniaczy operacyjnych jest zaprojektowanie uniwersalnej płytki drukowanej, na której można łatwo zmontować większość standardowych konfiguracji wzmacniaczy operacyjnych.

Zdobycie płytki drukowanej

Kiedyś, w sieci RS dostępna była płytka do testowania wzmacniaczy operacyjnych, pokazana na **rysunku 1**. Firma Maplin również produkowała takie płytki. Przez lata używałem wielu wycofanych już



Rysunek 8. Na płytce przewidziano dużo wolnego miejsca dla rezystora wyjściowego R12. Pozwala to na zastosowanie cewki izolującej wyjście wzmacniacza od obciążenia pojemnościowego



Rysunek 9. Wzmacniacz odwracający, standardowa konfiguracja wzmacniacza operacyjnego

z handlu układów, na przykład stosowanych we wzmacniaczu testowym EPE i kompresorze Upwards Compressor pokazanym na **rysunku 2**. Ponieważ tak przydatna płytką nie jest już dostępna w handlu, nadszedł czas aby stworzyć jej zamiennik, a przy okazji zastąpić elementy przewlekane z końca

lat siedemdziesiątych współczesnymi elementami SMD.

Płytki RS były stosowane w oprzyrządowaniu przemysłowym. Kiedy używałem ich w układach akustycznych, często zdarzało się, że duże kondensatory zwisały niebezpiecznie wokół płytki (**rysunek 3**). Dlatego w moim

nowym projekcie konieczne było uwzględnienie dobrego zamocowania dużych elementów. Z drugiej strony, niektóre wymagania dotyczące oprzyrządowania przemysłowego, takie jak duży zakres regulacji przesunięcia napięcia wyjściowego, można było pominąć. Potrzebna była również szyna do polaryzacji wejść podczas pracy z jednoprzewodowym zasilaniem. Ponadto jeden z najpopularniejszych wzmacniaczy operacyjnych stosowanych w układach akustycznych, czyli NE5534, wymaga czasami użycia kondensatora kompensacyjnego, włączonego między piny 5 i 8.

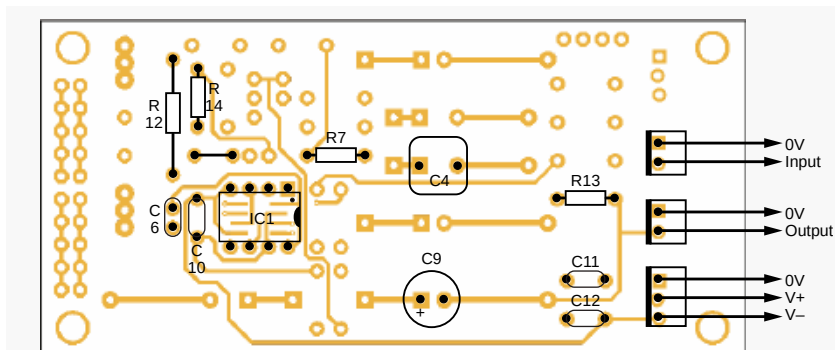
Układy zwielokrotnione

Niektóre z najlepszych akustycznych wzmacniaczy operacyjnych, takie jak LM4562, są dostępne tylko jako układy podwójne. Ogólnie rzecz biorąc, układy podwójne są bardziej popularne od pojedynczych, więc w przyszłości będzie również potrzebna inna płytka dwukanałowa, której projekt pojawi się później. Czterokanałowe wzmacniacze operacyjne, takie jak TL074, są rzadko stosowane w układach akustycznych, ponieważ trudno jest rozmieścić wokół nich duże kondensatory. Poza układem LM837 istnieje niewiele niskoszumowych wzmacniaczy czterokanałowych, więc nie będziemy się nimi zajmować. Uważam jednak, że przydatna byłaby płytka adaptera, zawierająca czterokanałowy układ wzmacniający, zbudowany na wzmacniaczach podwójnych. Pozwoliłoby to zastąpić kilka szumiących układów poczwórnych w niektórych z moich urządzeń z układami NE5532. Czy ktoś zaprojektował coś takiego? Jeśli będzie wystarczające zapotrzebowanie, wykonam płytkę drukowaną z konwerterem, z dwoma 8-pinowymi układami podwójnymi zastępującym jeden 14-pinowy układ poczwórny.

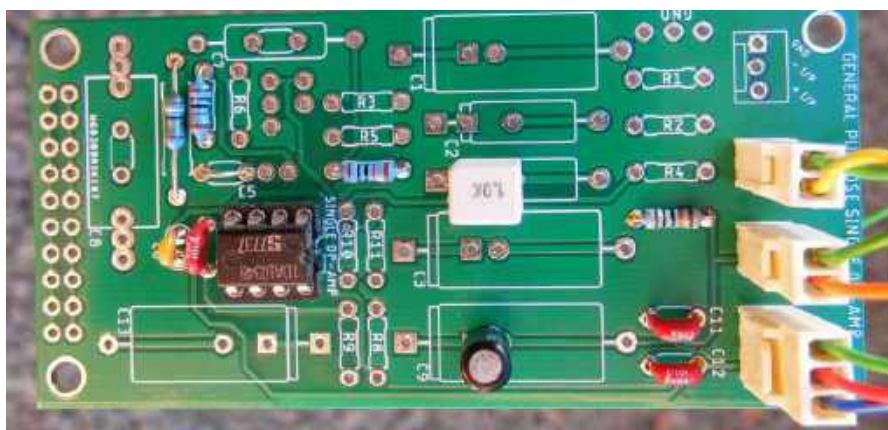
Płytką z dobrze rozmieszczonymi elementami

Schemat ideowy płytki pokazano na **rysunku 4**. Jest ona przystosowana do zastosowań akustycznych, ale z pewnością jej rola nie ogranicza się tylko do tego. W konkretnym układzie zostaną wstawione tylko niektóre elementy. Płytką ze wszystkimi możliwymi elementami pokazana jest na **rysunku 5**.

Wszystkie miejsca przeznaczone dla dużych kondensatorów mają dodatkowe



Rysunek 10. Rozmieszczenie elementów wzmacniacza odwracającego



Rysunek 11. Zmontowana płytka – zastosowano wzmacniacz operacyjny TDA1034, którego pierwowzorem jest NE5534 firmy Philips. Układ pochodzi z roku 1977 i był używany w konsoli mikserskiej grupy Pink Floyd, wyprodukowanej przez firmę Midas dla Britannia Row. Nadal działa dobrze, mimo że był wielokrotnie instalowany

punkty lutownicze, dostosowane do elementów o różnych obrysach. Umożliwia to stosowanie elementów stojących lub osiowych oraz kondensatorów połączonych równolegle. Po lewej stronie płytki znajdują się cztery grupy połączonych ze sobą punktów lutowniczych, służących do montażu dodatkowych elementów podczas eksperymentowania. Przykładowo, te dodatkowe punkty ułatwiły stworzenie elementu C8, składającego się z dwóch kondensatorów połączonych równolegle (rysunek 6).

Funkcje poszczególnych elementów

Zwykle w przypadku tego typu płytek elementy są opatrzone adnotacjami, zgodnie z ich funkcją w układzie. Na przykład na płytce RS rezystor ujemnego sprzężenia zwrotnego nazwano Rb. Aby uprościć zadanie projektantowi i konstruktorowi płytki drukowanej, postanowiłem pozostać przy takiej konwencji numeracji.

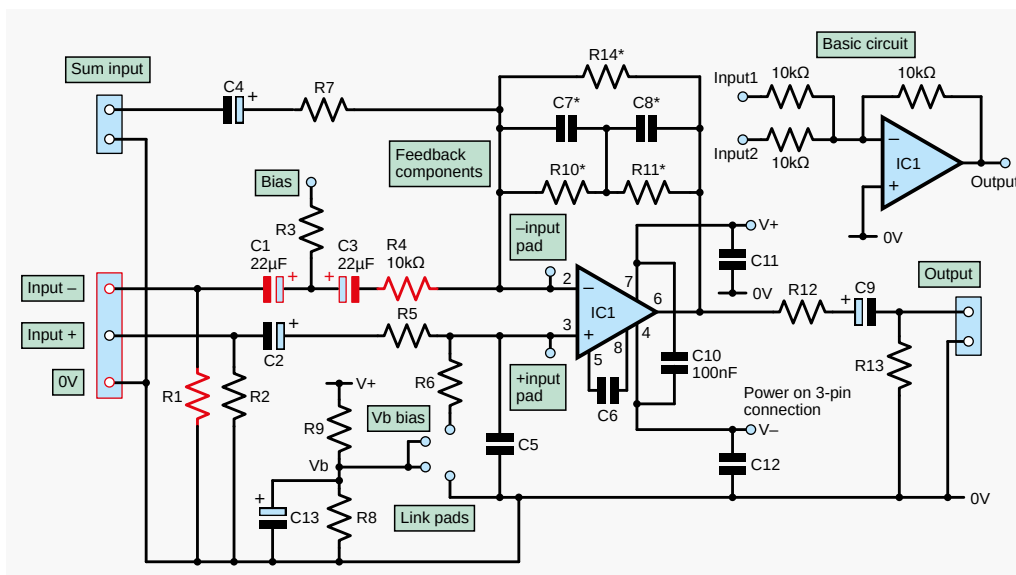
Poniżej znajduje się lista elementów składowych płytki, z opisem ich możliwych funkcji. Wszystko stanie się jasne, gdy zostaną omówione konkretne układy, w których wiele miejsc na elementy pozostanie pustych.

R1, R2 i R13

Rezystory te stanowią ważny dodatek w układach akustycznych. Są to rezystory wyrównujące napięcia na kondensatorach, przez co eliminowane są głośne trzaski powstające podczas podłączania wejść i wyjść układu do innych urządzeń. Rezystory te utrzymują wejściowe i wyjściowe wyprowadzenia kondensatorów sprzęgających na poziomie 0 V. Typowe wartości wynoszą kilkadziesiąt kΩ w przypadku elektrolitycznych kondensatorów sprzęgających, lub około kilkaset kΩ w przypadku kondensatorów foliowych i tantalowych (ze względu na ich mniejszą upływność).

C1, C2 i C3

Są to kondensatory sprzęgające, umieszczone na wejściach wzmacniacza operacyjnego. C1 i C3 są podłączone do wejścia odwracającego. W celu uzyskania niskiego poziomu zniekształceń mogą stanowić układ symulujący kondensator niepolaryzowany. Zniekształcenia powodowane przez kondensatory można dodatkowo zmniejszyć,



Rysunek 12. Wzmacniacz sumujący lub miksujący z dodatkowym wejściem – kondensator bipolarny (złożony z C1 i C3) jest zoptymalizowany pod kątem dobrego przenoszenia niskich częstotliwości. Uwaga: czerwone elementy pokazane na rysunku stanowią uzupełnienie w stosunku do elementów pokazanych na rysunku 9

polaryzując połączone ze sobą wyprowadzenia dodatnie elementów C1 i C3 napięciem 5 V. Odbywa się to za pomocą rezystora R3 połączonego z punktem Vb. Kondensatory C1 i C3 mogą również pełnić rolę elementów sprzężenia zwrotnego dla konfiguracji nieodwracającej. W takim przypadku wyprowadzenia rezystora R1 powinny być ze sobą zwarte. Kondensator C2 jest podłączony do wejścia nieodwracającego.

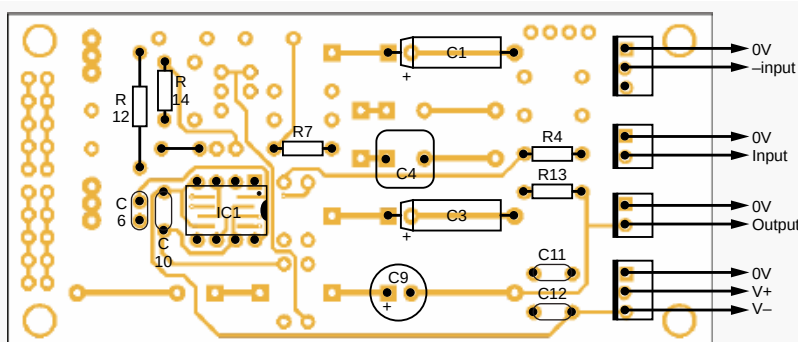
R4

Jest to rezystor wejściowy dla wzmacniacza odwracającego, zwykle oznaczany jako „Rin”. Może to być również rezystor wejściowy dla wejścia „minusowego” lub „zimnego” w terminologii akustycznej, we wzmacniaczu różnicowym. We wzmacniaczach odwracających, rezystor R6 jest zwykle połączony z masą układu za pomocą odpowiedniej zwory. W układach akustycznych często ma on niską wartość, np. 560 Ω, aby uzyskać niski poziom szumów, z jednoczesnym ograniczeniem

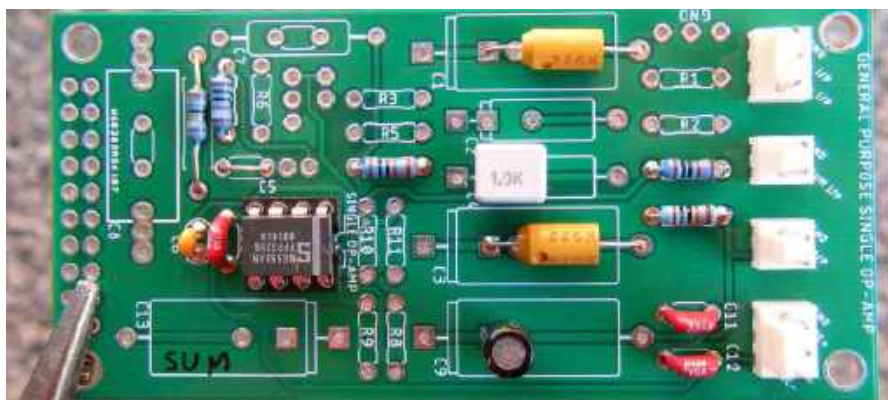
natężenia prądu wejściowego w przypadku awarii, takich jak nieprawidłowe podłączenie zasilania. W oprzyrządowaniu przemysłowym jest to zazwyczaj ta sama wartość, co rezystancja w pętli sprzężenia zwrotnego, gdyż to zapewnia najniższe przesunięcie dla prądu stałego. W przypadku zasilania układu z pojedynczego źródła napięcia, rezystor R6 jest podłączany za pomocą innej zwory do dzielnika dostarczającego połowę napięcia zasilającego, składającego się z rezystorów R8 i R9. Punkty lutownicze pozwalające na dokonanie wspomnianych połączeń są pokazane na rysunku 7.

R5

We wzmacniaczu różnicowym może to być rezystor wejściowy dla wejścia „plusowego” lub „gorącego” w terminologii akustycznej. W przypadku konfiguracji nieodwracającej rezystor R5 w połączeniu z kondensatorem C5 tworzą filtr dolnoprzepustowy, eliminujący zakłócenia o wielkiej częstotliwości.



Rysunek 13. Rozkład elementów we wzmacniaczu sumującym



Rysunek 14. Gotowy wzmacniacz sumujący

C4 i R7

Te elementy tworzą drugie wejście we wzmacniaczu odwracającym, gdy jest on używany jako sumator z wirtualną masą lub mieszać sygnałów.

R14

Jest to niezwykle ważny rezystor umieszczony w pętli ujemnego sprzężenia zwrotnego. Istnieje możliwość zastosowania bardziej złożonych układów sprzężenia zwrotnego, takich jak w przypadku korektora RIAA, składających się z elementów R10, R11, C7 i R8.

Na koniec omówmy kilka elementów zapewniających stabilność pracy układu. Jeśli potrzebny jest kondensator połączony równolegle z rezystorem sprzężenia zwrotnego R14, może to być C7 połączony szeregowo z C8. Typowe wartości to 33 pF do 100 pF. Rezystor R12 jest potrzebny do izolacji wyjścia wzmacniacza operacyjnego od obciążeń pojemnościowych, takich jak ekranowane kable. Zamiast tego, w celu uzyskania niższej

impedancji wyjściowej, można zastosować cewkę składającą się z 40 zwojów drutu, nawiniętego na rezystorze o mocy 1 W i rezystancji 39 Ω . Można także zastosować dławik o indukcyjności 10 μ H połączony równolegle z rezystorem (rysunek 8). Na stabilność pracy układu wpływają także kondensatory odsprężające C10, C11 i C13. Kondensator C6 jest kondensatorem kompensacyjnym dla wzmacniaczy operacyjnych NE5534, dla wzmocnień poniżej 5.

Lista komponentów (i ich funkcje)

Półprzewodniki

IC1 Pojedynczy wzmacniacz operacyjny, taki jak NE5534. Może to być 8-pinowy element w obudowie DIP do montażu przewlekane lub element w obudowie SOIC do montażu powierzchniowego.

Rezystory

Są to standardowe rezystory o mocy 0,25 W, zazwyczaj metalizowane. W układach

akustycznych należy stosować rezystory MRS25 o mocy 0,6 W i tolerancji 1%.

R1, R2, R13 rezystory wyrównujące stałe napięcia wejściowe i wyjściowe, o rezystancji od 22 k Ω do 100 k Ω .

R3 rezystor polaryzujący układ kondensatora bipolarnego lub zapewniający połączenie z masą w przypadku zastosowania dużego kondensatora w pętli sprzężenia zwrotnego, umieszczonego w polu C3.

R4, R5 rezystory wejściowe wzmacniacza operacyjnego, zazwyczaj o rezystancji od 1 k Ω do 100 k Ω .

R6 rezystor zerujący lub polaryzujący nieodwracające wejście wzmacniacza, zazwyczaj o rezystancji od 1 k Ω do 100 k Ω (lub znacznie wyższej, od 1 M Ω do 4,7 M Ω , dla wzmacniaczy operacyjnych z wejściem FET).

R7 rezystor tworzący wejście sumujące układu.

R8, R9 rezystory tworzące dzielnik dostarczający połowę napięcia zasilającego, mają równe wartości, od 10 k Ω do 100 k Ω każdy.

R10, R11 dodatkowe rezystory sprzężenia zwrotnego używane do tworzenia filtrów.

R14 główny rezystor w pętli sprzężenia zwrotnego, o wartości od zera do 220 k Ω (lub znacznie większej, od 1 M Ω do 4,7 M Ω , dla wzmacniaczy operacyjnych z wejściami FET).

R13 rezystor izolujący wyjście układu od obciążenia pojemnościowego, o rezystancji od 39 Ω do 600 Ω .

Kondensatory

C1, C2 wejściowe kondensatory sprzęgające, zazwyczaj o pojemności od 1 μ F do 22 μ F

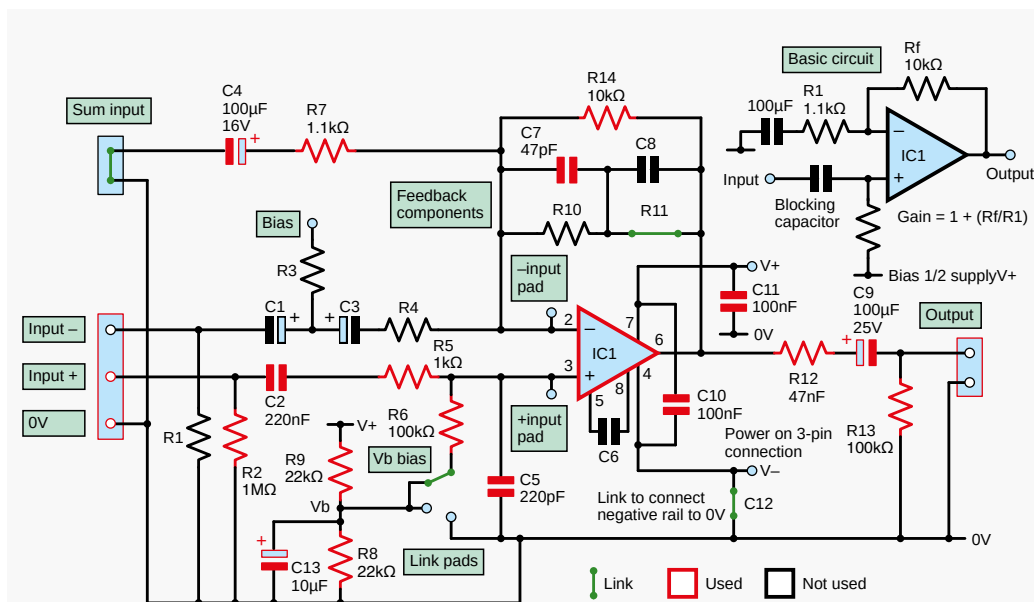
C3 drugi składnik kondensatora bipolarnego powstającego w połączeniu z C1 lub duży kondensator w pętli sprzężenia zwrotnego.

C4 dodatkowy, wejściowy kondensator sprzęgający w układzie sumującym lub mały kondensator w pętli sprzężenia zwrotnego.

C5 Składnik filtra eliminującego sygnały w.c.z., o pojemności od 47 pF do 470 pF.

C kondensator kompensujący, stosowany w układzie NE5534, o pojemności 22 pF dla wzmocnienia równego 1, zbędny dla wzmocnień powyżej 5, o pojemności 4,7 pF dla korektorów RIAA, rozstaw 2,5 mm.

C7, C8 kondensatory pracujące w pętli sprzężenia zwrotnego, pole C8 jest przeznaczone dla bardzo dużych, poliestrowych kondensatorów filtrujących.



Rysunek 15. Wzmacniacz nieodwracający o wzmocnieniu 10x i zasilaniu jednoprzewodowym

C10 kondensator odsprężający piny zasilające układu scalonego, ceramiczny, o pojemności 0,1 μF , rozstaw 5 mm

C11, C12 kondensatory odsprężające szyny zasilające układu, ceramiczne o pojemności 0,1 μF lub elektrolityczne o pojemności do 10 μF , rozstaw 5 mm.

Złącza

Złącza PCB Molex o rozstawie 0,1 cala lub ich odpowiedniki:

złącze dwustykowe, proste JYK P2500-02, 2 szt.

złącze trójstykowe, proste JYK P2500-03, 2 szt.

Kody do szybkiego zamówienia 22-0950 i 22-0955

Klasyczne konfiguracje układowe

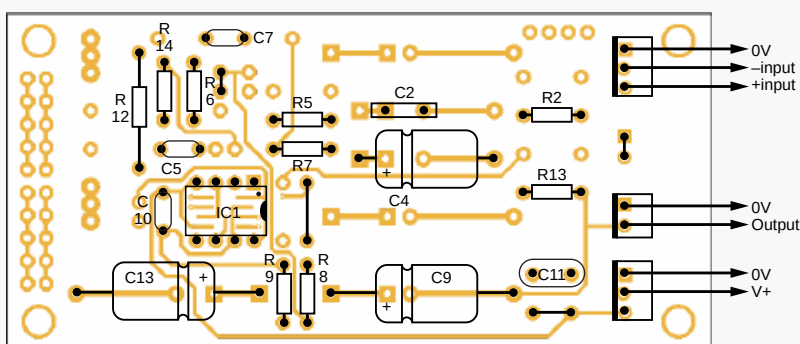
Oto kilka przykładowych układów, które można zbudować na opisywanej płytce. Jestem pewien, że Czytelnicy dostosują ją do wielu innych zastosowań – prosimy o przesłanie nam swoich pomysłów.

Wzmacniacz odwracający

Jest to najprostszy układ wzmacniacza, jaki można zbudować na tej płytce. Istnieją dwa sposoby sterowania wejść. Jednym z nich jest doprowadzenie sygnału do dwustykowego złącza Molex pokazanego na schemacie (rysunek 9) i na rysunku z rozkładem ścieżek (rysunek 10). Alternatywą jest użycie wejścia odwracającego, dostępnego na złączu trójstykowym i doprowadzenie sygnału wejściowego przez kondensator C1. Możliwe jest również użycie niepolaryzowanego kondensatora sprzęgającego, złożonego z elementów C1 i C3. W celu uzyskania możliwie niskich zniekształceń sygnału można ten układ spolaryzować, przez podanie napięcia zasilającego na rezystor R3. Jest to przydatne w przypadku układów o niskiej impedancji wejściowej, w których rezystor wejściowy ma małą rezystancję. Gotową płytę ze wzmacniaczem odwracającym pokazano na rysunku 11.

Wzmacniacz sumujący

Jest to wzmacniacz odwracający z dodatkowym rezystorem w układzie wejściowym. Wykorzystane są oba wejścia, dostępne na złączach J1 i J3. Wynikowy układ pokazano na rysunku 12. W obwodzie sprzężenia zwrotnego zastosowany został układ składający się z kondensatorów bipolarnych. Sygnał wyjściowy z filtra dolnoprzepustowego został zmieszany z sygnałem z korektora. Powstał w ten sposób filtr wycinający typu notch. W obwodzie wyjściowym wzmacniacza zastosowany został kondensator o wysokiej



Rysunek 16. Rozkład elementów we wzmacniaczu nieodwracającym – zwróć uwagę na połączenia w układzie polaryzacji



Rysunek 17. Kompletny wzmacniacz nieodwracający

pojemności, równej 100 μF , przez co uniknięto spadku wzmocnienia w zakresie niskich częstotliwości. Rozkład elementów jest pokazany na rysunku 13, zaś wygląd zmontowanej płytki na rysunku 14.

Wzmacniacz nieodwracający

Jest to jeden z najpopularniejszych układów zbudowanych w oparciu o wzmacniacz operacyjny, zapewniający wzmocnienie w zakresie od 1 do 1000 (0 dB do 60 dB) oraz wysoką impedancję wejściową. Rysunek 15 przedstawia układ o wzmocnieniu 10x.

Układ ten jest nieco bardziej skomplikowany od innych, ponieważ został przystosowany dla zasilania jedнопроводового (a nie typowego zasilania dwuprzewodowego). Osiągnięto to przez polaryzację wejścia nieodwracającego napięciem o połowę mniejszym od napięcia zasilającego. Rozkład elementów został przedstawiony na rysunku 16. W takiej konfiguracji jest budowany przedwzmacniacz mikrofonowy. Składa się on ze wzmacniacza nieodwracającego o regulowanym wzmocnieniu i transformatora podwyższającego napięcie sygnału wejściowego. Na rysunku 17 przedstawiony

jest wygląd gotowej płytki. By zmienić sposób zasilania na dwuprzewodowy należy pominąć dzielnik dostarczający napięcie polaryzujące i połączyć rezystor R6 z masą.

W następnym miesiącu

W części 2 zakończymy opis tego projektu, przedstawimy wzmacniacz różnicowy i wzmacniacz gramofonowy z korekcją RIAA, a także nietypowy, a za razem przydatny wzmacniacz operacyjny z kompensacją niskich częstotliwości. ■

Jake Rothman

Ten artykuł był wcześniej opublikowany na łamach „Practical Electronics”, grudzień 2022 (www.epemag3.com)

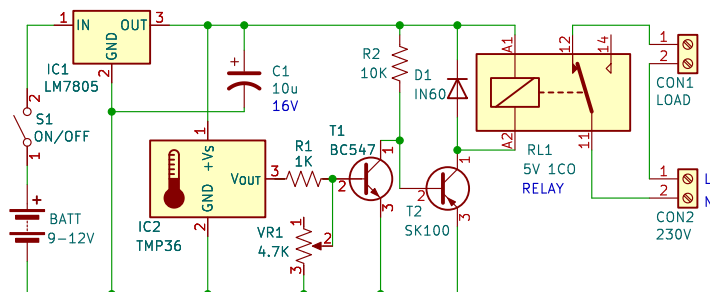
Prosty i niezawodny czujnik temperatury

Współczesna elektronika ewoluuje w kierunku bardzo rozbudowanych konstrukcji. Główny kierunek rozwoju to technika komputerowa z szerokim wsparciem informatycznym. Trudno dzisiaj o projekt który nie zawierałby mikroprocesora, mikrokontrolera lub wręcz nie angażowałby komputera z bogatym i zaawansowanym software-em.

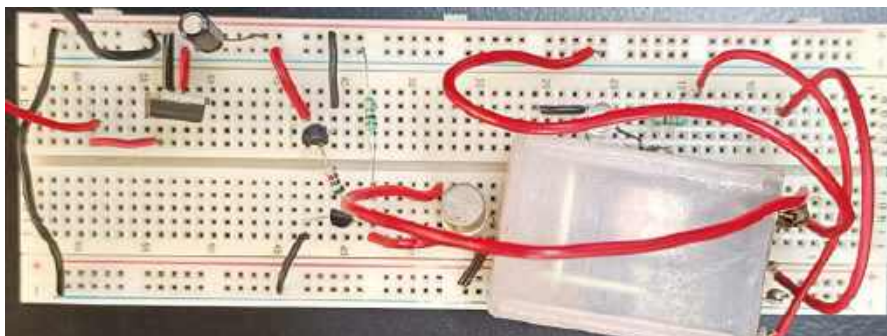
Mimo to, istnieje zapotrzebowanie na projekty o wiele prostsze, gdzie zamierzony cel można osiągnąć nie angażując skomplikowanego hardware-u i bez software-u. Jednym z obszarów otwartych dla prostej elektroniki jest pomiar temperatury. Trudno wymienić wszystkie urządzenia, gdzie pomiar temperatury jest niezbędny lub bardzo pożądanym. A są wśród nich i proste kuchenki czy podgrzewacze żywności, lodówki, elektryczne czajniki, kuchenki mikrofalowe lub mniej i bardziej skomplikowane systemy klimatyzacji. Jeśli potrzebny jest elektroniczny termometr, można polecić wiele czujników w oparciu o które można takowy zbudować. Jednym z nich i godnym polecenia jest układ scalony TMP36. Układ ten mierzy temperaturę w szerokim zakresie i zachowuje dobrą dokładność. Nic nie stoi na przeszkodzie, aby współpracował z mikroprocesorem, który dopiero uruchomi urządzenia wykonawcze. Bieżący projekt nie angażuje software-u. To samo zadanie wypełni prosty driver pośredniczący między termometrem i wykonawczym przełącznikiem. **Rysunek 1** pokazuje zdjęcie prototypu zmontowanego na płytce uniwersalnej. To obwód który przeszedł pomyślne testy w laboratorium naszej Redakcji „Electronics For You”.

Budowa układu i jego działanie

Schemat ideowy sprytnie zaprojektowanego układu pomiaru temperatury pokazano na **rysunku 2**. Wykorzystano tu niewielką liczbę elementów i zasadniczymi są:



Rysunek 2. Schemat ideowy układu

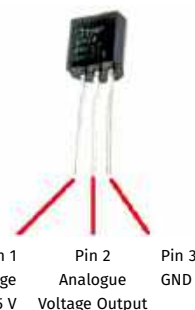


Rysunek 1. Prototyp poddany testom

5 V stabilizator napięcia LM7805 (IC1), czujnik temperatury TMP36 (IC2), 5-cio voltowy przełącznik (RL1), tranzystor NPN BC547 (T1), tranzystor PNP SK100 (T2) i dioda 1N60 (D1). Ponadto jest tylko kilka dyskretnych pasywnych elementów.

Sercem naszego projektu jest scalony sensor TMP36. Najistotniejsze cechy tego elementu są następujące: szeroki zakres napięcia zasilania (2,7 V do 5,5 V); zakres pomiaru temperatury od -40°C do $+150^{\circ}\text{C}$; liniowa charakterystyka ($y=mx+c$); nachylenie charakterystyki $20\text{ mV}/^{\circ}\text{C}$; niewielki prąd spoczynkowy, poniżej $50\text{ }\mu\text{A}$. TMP36 może pracować w warunkach zmienności temperatury do $3^{\circ}\text{C}/\text{s}$ gdy temperatura rośnie i do $6^{\circ}\text{C}/\text{s}$ gdy układ „stygnie”. **Rysunek 3** pokazuje kolejność wyprowadzeń czujnika TMP36 w obudowie TO92.

Zasada działania układu wg bieżącego projektu jest prosta. Czujnik TMP36 działa jak przetwornik temperatura-napięcie. Zatem,



Rysunek 3. Funkcje wyprowadzeń IC TMP36

sygnałem wyjściowym jest napięcie i jest to „odczyt” w zakresie miliwoltów. Kiedy podgrzejemy sensor TMP36, napięcie na bazie tranzystora T1 podnosi się i w pewnym momencie zostanie przekroczony próg przewodzenia złącza baza-emiter. Gdy oba

Wykaz elementów:

Półprzewodniki:

IC1: 7805 – stabilizator 5 V
IC2: TMP36 – czujnik temperatury
D1: dioda 1N60
T1: BC547 – tranzystor NPN
T2: SK100 – tranzystor PNP

Rezystory:

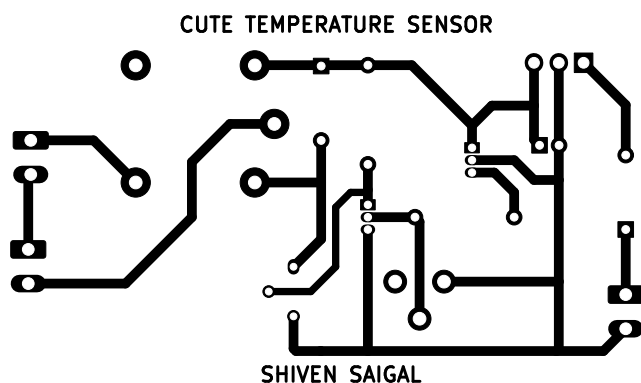
(wszystkie 0,25 W/±5%)
R1: 1 kΩ
R2: 10 kΩ
VR1: 4,7 kΩ – potencjometr wieloobrotowy

Kondensatory:

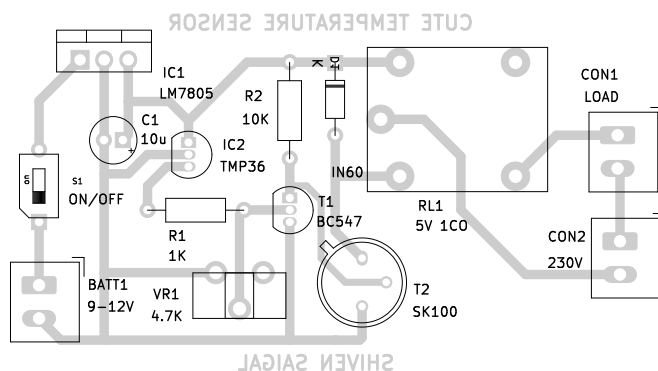
C1: 10 μF/16 V – elektrolityczny

Pozostałe:

CON1, CON2: złącze 2-pinowe
RL1: przełącznik 5V NO/NC
S1: przełącznik ON/OFF
BATT1: bateria 9–12 V



Rysunek 4. Projekt druku PCB



Rysunek 5. Schemat montażowy elementów na PCB

tranzystory T1 i T2 przewodzą prąd w obwodzie kolektor-emiter, zostanie włączony przełącznik którego styki uruchomią układ wykonawczy w danej aplikacji (bliżej nieokreślony przez konstrukcję samego czujnika będącego tematem projektu). Jako driver cewki przełącznika wykorzystano parę tranzystorów npn-pnp w tzw. połączeniu Sziklaiego. Układ taki oferuje podobne wzmocnienie do tradycyjnego połączenia Darlingtona, aczkolwiek w bieżącym zastosowaniu ma nad nim przewagę. Proponowany „temperature sensor” pracuje w zakresie niskich napięć, a driver w postaci połączenia Darlingtona potrzebowałby napięcia na bazie w okolicy 1,22 V. Dla włączenia pary Sziklaiego wystarczy połowa tego napięcia, jak w przypadku pojedynczego tranzystora (ok. 0,61 V).

Kalibracja układu

Czujnik temperatury trzeba skalibrować w zależności od pożądanej temperatury pracy. Tu przewidziano wieloobrotowy potencjometr, rezygnując z rozbudowanego systemu zawierającego mikrokontroler. Procedura kalibracji jest tu wyjątkowo prosta. Należy rozrządzić czujnik do temperatury przy której chcemy aby przełącznik, np. odłączył obciążenie od zasilania. Teraz należy kręcić suwakiem potencjometru do momentu, aż przełącznik przyciągnie kotwicę a tym samym odłączy obciążenie od zasilania (wykorzystanie styku „normalnie załączonego” NC). Gdy temperatura opadnie, cewka przełącznika zwolni kotwicę a ta ponownie zewrze styk NC przełącznika, podłączając tym samym obciążenie do zasilania. Zazwyczaj wykorzystuje się styk „NO” do sterowania obciążeniem, nic jednak nie stoi na przeszkodzie, by użyć (jak tu) styku „NC” (normalnie zwarte) właśnie po to, by rozewrzeć sterowany obwód w momencie załączenia się przełącznika. Obwód wykonawczy może załączać lub rozłączać ładowarkę, kuchenkę, lutownicę lub inny gadżet. Należy sprawdzić,

czy po obniżeniu temperatury poniżej uznanej za krytyczną, obciążenie zostanie ponownie zasilone. Sensor TMP36 może być rozgrzewany z szybkością 3°C/s i schładzany 6°C/s. Zwykle procesy termiczne nie pozwalają na szybką dynamikę zmian temperatury i większość ciał stygnie wolno o ile nie jest uruchomiony system szybkiego chłodzenia. Należy oczekiwać, że nasz czujnik temperatury powinien pracować stabilnie bez oscylacji wokół temperatury zadanej.

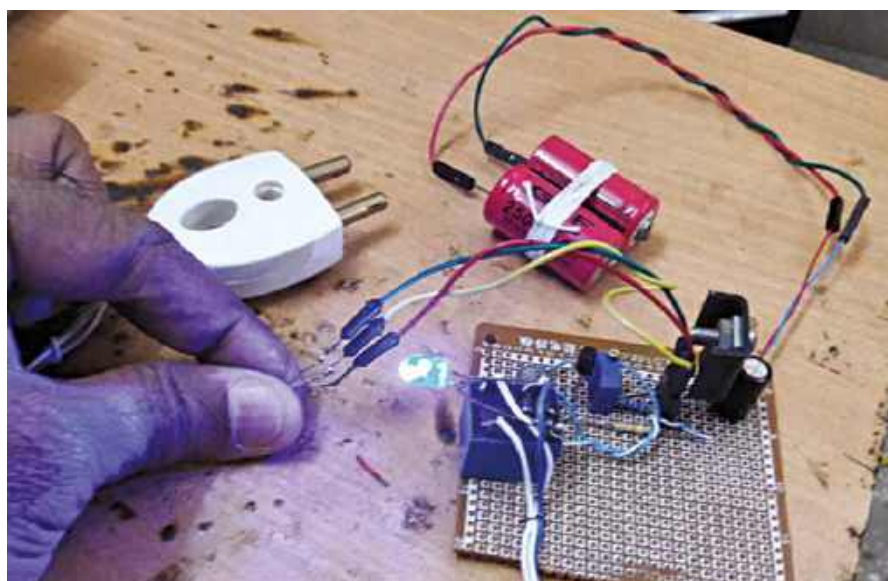
Konstrukcja układu i testowanie jego działania

Na potrzeby projektu przygotowano płytkę PCB z jednostronnym drukiem. **Rysunek 4** powinien być w skali 1:1, zaś na **rysunku 5** pokazano ułożenie elementów na PCB.

Po zmontowaniu układu należy go zamknąć w odpowiednio przygotowanej obudowie. Alternatywnie, układ można zmontować na płytce uniwersalnej i również przygotować stosowną obudowę. **Rysunek 6** pokazuje takie rozwiązanie wykonane przez autora.

W obudowie należy przewidzieć złącza zarówno na wejściu jak i wyjściu, szczególnie jeśli zasilanie bateryjne będziemy chcieli zastąpić transformatorem obniżającym napięcie z sieci 230 V AC. Ten prosty układ będący tematem bieżącego projektu może mieć wiele zastosowań. Możemy tu przytoczyć kilka przykładów.

- Kontrola temperatury akumulatorów Ni-Mh lub litowo-jonowych podczas ładowania. Baterie podczas ładowania rozgrzewają się i jako bezpieczny przedział przyjmuje się od 45 do 50°C dla Ni-Mh oraz do 45°C dla Li-Ion i Li-Po. Kontrola temperatury i ewentualne przerwanie procesu ładowania wydłuży żywotność akumulatorów ponad to, co oferują zwykle ładowarki.
- Kontrola temperatury grota lutownicy. Tu należałoby wykorzystać styki NC. Równocześnie należy zadbać aby czujnik TMP36 śledził temperaturę grota, ale zakres temperatury czujnika jest ograniczony do +150°C. Zatem,



Rysunek 6. Prototyp czujnika temperatury wykonany przez autora

powinien być on nieco oddalony od grota lutownicy.

- Kontrola wielkości płomienia palnika gazowego, gdy podczas grzania wody zaczyna się ona gotować. W takim zastosowaniu należy wykorzystać styki „NO” i za ich pomocą sterować serwowymechanizmem regulującym wielkość płomienia. Termometr TMP36 najlepiej żeby miał bezpośredni kontakt z grzaną/gotowaną wodą. W tym przypadku należy zadbać o dobry kontakt termiczny, ale dla samego czujnika trzeba przygotować obudowę wodoszczelną.

W takim zastosowaniu najistotniejszym fragmentem projektu może być właśnie mocowanie czujnika. Użycie styków „NO” lub NC, to zależy od budowy serwowymechanizmu który reguluje zaworem otwierającym/zamykającym dopływ gazu. Musi być oczywiście zachowane ujemne sprzężenie zwrotne, ale problem ze stabilnością może być poważny. Kilka słów na ten temat zamieszczamy w komentarzu na końcu opracowania – przy pracy redakcji.

W świecie zdominowanym przez zaawansowaną technologię, ten projekt przypomina nam, że prostota też „daje moc”. Często właśnie proste, ale pomysłowe rozwiązania otwierają drzwi do niezliczonych możliwości poprawy naszego codziennego życia. ■

Somnath Bera

Ten artykuł był wcześniej opublikowany na łamach „EFY”, sierpień 2023 (efymag.com)

Układ jest bardzo prosty, mimo to jest tu kilka błędów. Autor chciał uniknąć w stopniu drivera połączenia dwóch tranzystorów w układzie Darlingtona, zastępując go parą pnp-npn Sziklaiego. Faktycznie, jeden tranzystor NPN nie zdałby tu egzaminu, a Darlington też nie byłby wskazany. Pomijając błąd, iż w tranzystorze T2 zamieniono emiter z kolektorem, to obwód Sziklaiego też nie jest dobry. Taki obwód nigdy się nie nasyci i w stanie włączenia napięcie między kolektorem i emitern T2 pozostanie na poziomie około 1 V. Przy zasilaniu 5 V to może być dużo. 5 V przełącznik przy czterech woltach powinien włączyć, ale w przypadku dużego przełącznika, cewce o niskiej rezystancji, tranzystor T2 może się trochę grzać. Ale to też nie jest największy mankament. Tutaj, komparatorem napięcia jest złącze baza-emiter tranzystora T1 i z tym się wiąże szereg problemów. To nie jest dobre rozwiązanie. Pierwsze, co rzuca się w oczy to fakt, iż to mało precyzyjny próg i też zależny

od temperatury. Charakterystyka wyjściowa TMP36 ma nachylenie 10 mV/°C, a napięcie złącza baza-emiter (przy stałym prądzie) obniża się o około 2,5 mV ze wzrostem temperatury o 1°C. Na dobrą sprawę układ scalony TMP36 działa w oparciu o tę samą zależność charakterystyki złącza w tranzystorze bipolarnym od jego temperatury. Ta zależność jest tylko odpowiednio wzmocniona i skaliowana. Jeśli IC2 i T1 są w jednakowej temperaturze, to zmienność U_{be} z temperaturą można wkalibrować w charakterystykę całości i nie musiałoby to być problemem. Można jednak bezpiecznie stwierdzić, iż trudno o gorszy komparator aniżeli bazowanie na charakterystyce złącza baza-emiter tranzystora T1. Jednak, to też nie jest największy szkopuł. Termometr TMP36 ma liniową charakterystykę, przyzwoitą dokładność i potrafi mierzyć temperaturę w szerokim zakresie od -40°C do +150°C. Ale nie tu! Można by argumentować, że szeroki zakres pomiaru nie jest (w pewnej aplikacji) potrzebny. Ale, skoro autor wybrał wersję TMP36, to chyba na szerokim zakresie pracy powinno mu zależeć. TMP37 ma ten zakres węższy, a charakteryzuje się dwukrotnie większym nachyleniem charakterystyki. I dwie rzeczy są tu kluczowe. Charakterystyka napięciowa wyjścia (względem temperatury) i obciążalność prądowa tego wyjścia. Charakterystyka $U=f(T)$ jest liniowa, ze stałym nachyleniem +10 mV/°C w pełnym zakresie pomiarowym i z przesunięciem zera (względem 0°C) +0,5 V (i w tym zakresie są główne różnice między TMP36 i termometrami pokrewnymi TMP35 i 37). Warto też dodać, iż charakterystyka ta jest praktycznie stała w pełnym zakresie zasilania układu scalonego.

Dla włączenia przełącznika, należy oczekiwać napięcia 0,6 V na bazie T1. Wyjście termometru da napięcie 0,6 V przy temperaturze około +10°C. Zatem pomiar niższych temperatur jest absolutnie niedostępny. Uwzględniając dzielnik rezystancyjny na wyjściu (i abstrahując na razie od obciążalności prądowej), nawet przy potencjometrze VR1 ustawionym na maksimum, należy tu wprowadzić poprawkę do napięcia z 0,6 V do 0,73 V, a na skali temperatury to już około 25°C. Zatem nie ma co marzyć o pomiarze temperatur niższych (a sam termometr ma zakres od -40°C). Powiedzmy, że autorowi projektu ten przedział nie był potrzebny. Ale, w drugą stronę jest jeszcze gorzej! Przy temperaturze 125°C termometr da napięcie około 1,7 V. Ale, obciążalność wyjścia to zaledwie 50 µA. Prąd bazy T1 zależy od zastosowanego przełącznika i wzmocnienia prądowego h₂₁ obu tranzystorów.

Można szacować, że w pięćdziesięciu mikroamperach się zmieści. Ale R1 o wartości 1 kΩ i VR1 do 4,7 kΩ stanowią poważny problem. Aby przestawić punkt pracy na wyższą temperaturę, trzeba VR1 skrócić w kierunku niższej rezystancji. Jak by to było np. w 125°C? Wyjście 3 IC1 da napięcie 1,7 V, a na bazie T1 chcemy +0,6 V. Nietrudno policzyć jaką rezystancję należy ustawić na VR1. Ale policzmy prąd czerpany z wyjścia IC2. 1,7 V - 0,6 V = 1,1 V. Podzielić przez 1 kΩ = 1,1 mA. Katalog podaje, iż prąd zwarcia wyjścia jest na poziomie 250 µA! W takim razie, w jakim zakresie możemy oczekiwać, że układ będzie działał? 50 µA × 1 kΩ = 50 mV. I to należałoby dodać do oczekiwanego 0,6 V na bazie T1. Ale jest jeszcze gorzej! Już przy napięciu 600 mV na bazie T1 i przy VR1 skróconym na max, prąd płynący przez ten potencjometr wyniesie 127 µA. Do tego należałoby dodać coś dla prądu bazy. Termometr po prostu takiego prądu nie da! I termostat ten nie będzie działał w żadnym zakresie! To w zasadzie zwalnia z poszukiwania dalszych błędów. A są nimi i brak histerezy pożądanej w każdym termostacie i problemy zysterowaniem przełącznika. I możliwość niestabilności, oscylacji, gdyby układ w ogóle działał. Istniałby także, choć wąski przedział temperatury w której napięcie na cewce przełącznika byłoby pośrednie między włączeniem i wyłączeniem. Choć minimalne dodatnie sprzężenie zwrotne należałoby dodać, aby takiego stanu uniknąć. Jako atut, autor reklamuje swój projekt, że jest prosty. Ale jest „za prosty” na to, aby poprawnie działał!

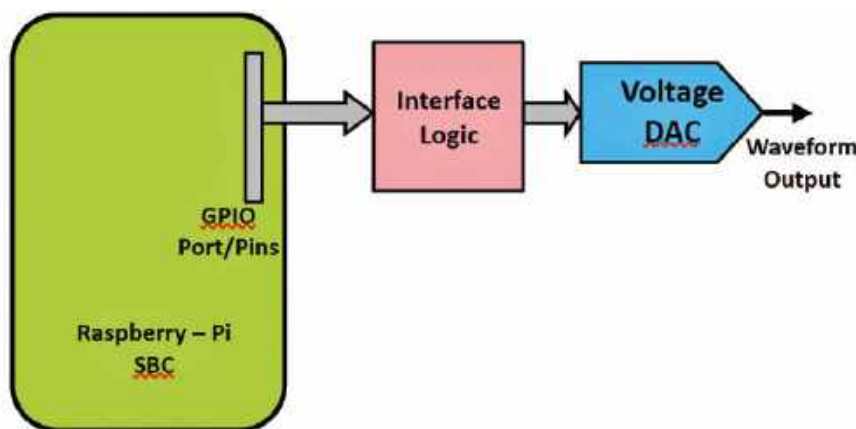
Wymienione błędy i bolączki układu z rysunku 2 nie jest trudno poprawić, i układ taki może być faktycznie użyteczny w wielu zastosowaniach. Odnosząc się do tekstu trudno nie odnieść wrażenia, że większość opisu to „masło maślane”. Autor przytacza też kilka parametrów elementu TMP36 odczytanych z katalogu. Ale np. błędnie zinterpretował dane Ramp-Up Rate 3°C/s i Ramp-Down Rate -6°C/s.

To parametry cytowane jako dopuszczalne w procesie lutowania układu scalonego. Nie są to parametry które zabezpieczą układ pracujący ze sprzężeniem zwrotnym przed ewentualnymi oscylacjami. To może być istotny problem wart odrębnego omówienia. Tu powiemy tylko, iż regulacja dwupołożeniowa, kiedy system „chodzi po histerezie” jest w miarę łatwa do ogarnięcia. Poważne zagrożenie niestabilnością jest wtedy, gdy oczekujemy szybkiej „odpowiedzi systemu” i regulator współpracuje z serwowymechanizmem, co autor podaje w ostatnim przykładzie (regulacji palnika gazowego).

Generator funkcyjny na Raspberry Pi

Generatory funkcyjne należą do przyrządów pomiarowych powszechnie używanych w pracowniach elektro-
nicznych. Wytwarzają przebiegi okresowe o różnych kształtach, częstotliwościach i amplitudach, przydatne
do testowania szerokiej klasy układów. I chociaż prezentowany tu układ ustępuje parametrami przyrządom
wysokiej klasy, tym niemniej doskonale się sprawdzi w każdej pracowni jako generator typowych przebie-
gów – sinusoidalnego, prostokątnego, trójkątnego i piłokształtnego. A ponieważ jest sterowany w pełni
programowo, oferuje dodatkową funkcję „generatora przebiegów arbitralnych”, umożliwiając użytkownikowi
definiowanie i wytwarzanie przebiegów o nietypowym kształcie.

Układ wykorzystuje popularny sterownik Raspberry Pi SBC oraz kilka innych łatwo dostępnych i niedrogich elementów. Zbuduje go bez problemu zarówno hobbysta jak i profesjonalista. Oprogramowanie zostało napisane w języku Python dla Raspberry Pi. **Rysunek 1** przedstawia schemat blokowy generatora. Elementy są wymienione w tabeli 1.



Rysunek 1. Schemat blokowy generatora przebiegów

Układ i działanie

Schemat generatora funkcyjnego przedstawiono na **rysunku 2**. Jest on zbudowany z wykorzystaniem mikrokomputera Raspberry Pi, 8-bitowego szeregowo-równoległego rejestru przesuwającego 74AHC595 i kilku innych elementów.

Działanie układu jest proste. Jak widać na schemacie z rysunku 2, centralną częścią generatora jest przetwornik cyfrowo-analogowy w formie drabinki R-2R z wyjściem napięciowym, składający się z rezystorów R2... R17. Przetwornik dostaje na wejściu słowo cyfrowe i zamienia je bezpośrednio na napięcie wyjściowe. Przetwornik cyfrowo-analogowy

z wyjściem prądowym wymagałby zastosowania wzmacniacza operacyjnego w celu konwersji prądu na napięcie. Zamiast tego został użyty prosty układ napięciowy, nie wymagający dodatkowego zasilania. Szczegóły dotyczące pracy i działania przetwornika cyfrowo-analogowego R-2R można znaleźć np. w numerze „Electronics For You”

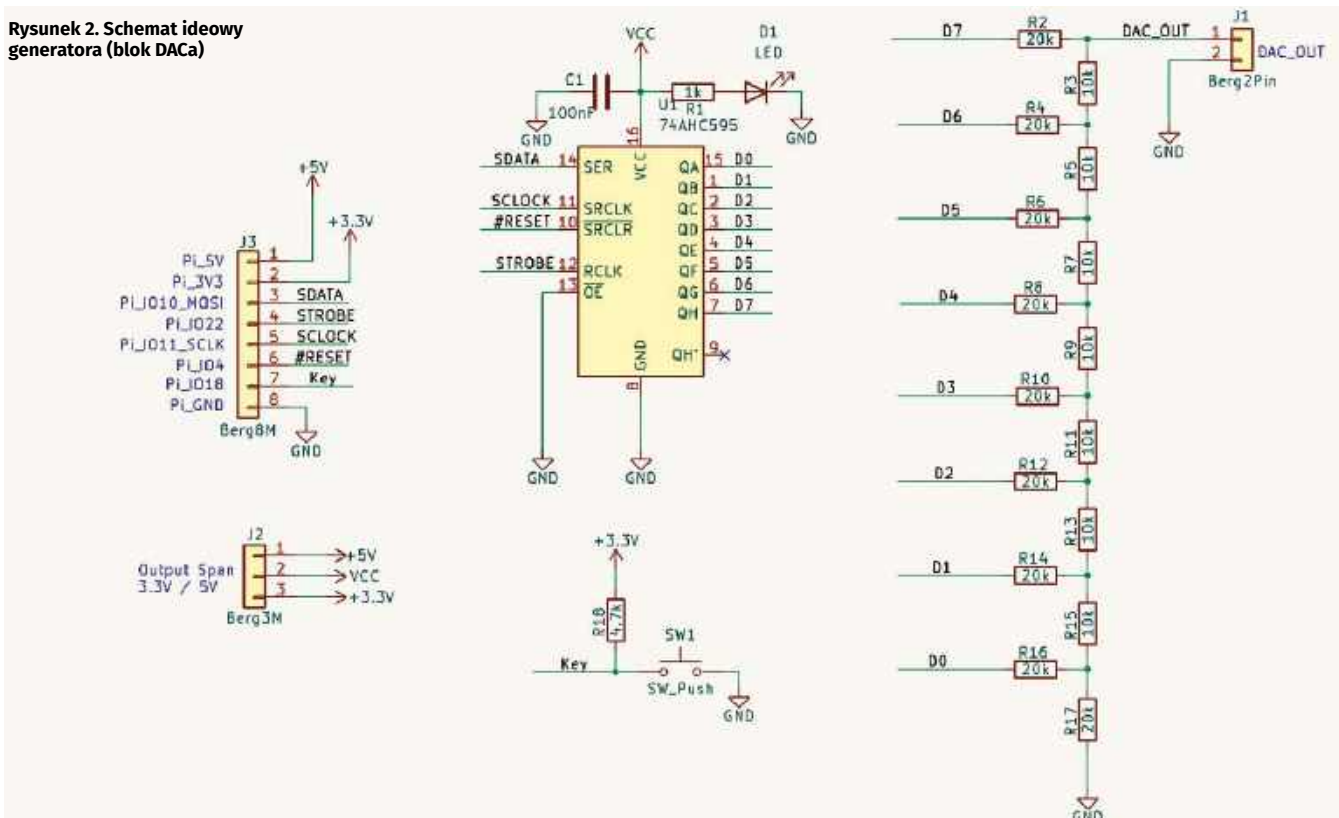
z października 2000 r. w artykule „R-2R DAC Based Waveform Generator Using PIC16C84 Microcontroller”. Wyjście przetwornika cyfrowo-analogowego wyprowadzono na złączce J1, oznaczone jako „DAC_OUT”.

Słowa cyfrowe dla DAC-a pochodzą z układu scalonego 74AHC595: 8-bitowego szeregowo-równoległego rejestru przesuwającego.

Tabela 1. Lista elementów

nazwa, symbol	wartość	komentarz
moduł Raspberry Pi	–	
74AHC595	IC	8-bitowy rejestr przesuwający
C1	100 nF	kondensator ceramiczny
D1	LED	5 mm, czerwony lub zielony
J1	2-stykowe złącze męskie	złącze Berg
J2	3-pinowe złącze męskie	złącze Berg (opcjonalne)
J3	8-stykowe złącze męskie	złącze interfejsu Raspberry Pi
R1	1 kΩ	rezystor 0,25 W, 5%, węglowy
R3, R5, R7, R9, R11, R13, R15	10 kΩ	rezystor 0,25 W, 1%, metalizowany
R2, R4, R6, R8, R10, R12, R14, R16, R17	20 kΩ	rezystor 0,25 W, 1%, metalizowany
R18	4,7 kΩ	rezystor podciągający; opcjonalny; patrz tekst
SW1	przycisk 6 mm	opcjonalny; patrz tekst
kable połączeniowe	10	kabelki żeńsko-żeńskie

Rysunek 2. Schemat ideowy generatora (blok DACa)



Przypis redaktora: układ ten, szeroko stosowany w technice mikroprocesorowej, składa się z rejestru przesuwającego z wejściem szeregowym oraz wyjściowego rejestru równoległego. Dane z Raspberry Pi są wysyłane przez wbudowany interfejs szeregowy SPI. Interfejs używa linii zegara (SCLK) na pinie GPIO 11 oraz linii danych (SDATA) na pinie GPIO 10, który jest „wyjściem układu nadrzędnego i wejściem podrzędnego” (MOSI). Oprócz tych dwóch głównych sygnałów transmisji szeregowej, do obsługi 74AHC595 potrzebne są jeszcze dwa sygnały. Jeden z nich to „strobe”, który szeregowo dane wprowadzone do rejestru przesuwającego zatrzymuje w rejestrze wyjściowym. Dane te (linie D7...D0) stanowią 8-bitowe słowo wejściowe przetwornika DAC. Sygnał strobujący jest wytwarzany programowo na pinie GPIO 22. Drugi sygnał to „reset”, który zeruje rejestr przesuwający w 74AHC595. „Reset” jest wytwarzany na pinie GPIO 4 i aktywowany tylko na początku.

Złącze J2 służy do wyboru napięcia zasilania rejestru przesuwającego: 3,3 V lub 5 V. Napięcie to określa poziomy na liniach D7...D0, a tym samym skalę napięcia przetwornika („zakres wyjściowy”), czyli maksymalną wartość napięcia wyjściowego (dla słowa 0xFF = 25510).

Przypis redaktora: karta katalogowa układu 74AHC595 zaleca, by przy zasilaniu

5 V napięcie wysokiego poziomu logicznego na jego wejściach wynosiło co najmniej 3,5 V. Tymczasem wyjścia Raspberry Pi mają w stanie wysokim napięcie odrobinę niższe: 3,3 V. Jeśli planowane jest zasilanie rejestru z 5 V, zaleca się użyć jego wersji dostosowanej do niższego napięcia logicznego jedyńki – 74AHCT595 lub 74HCT595.

Jest też wejście cyfrowe „Key” z przyciskiem (pin 18 GPIO w Raspberry Pi). Przycisk można wykorzystać

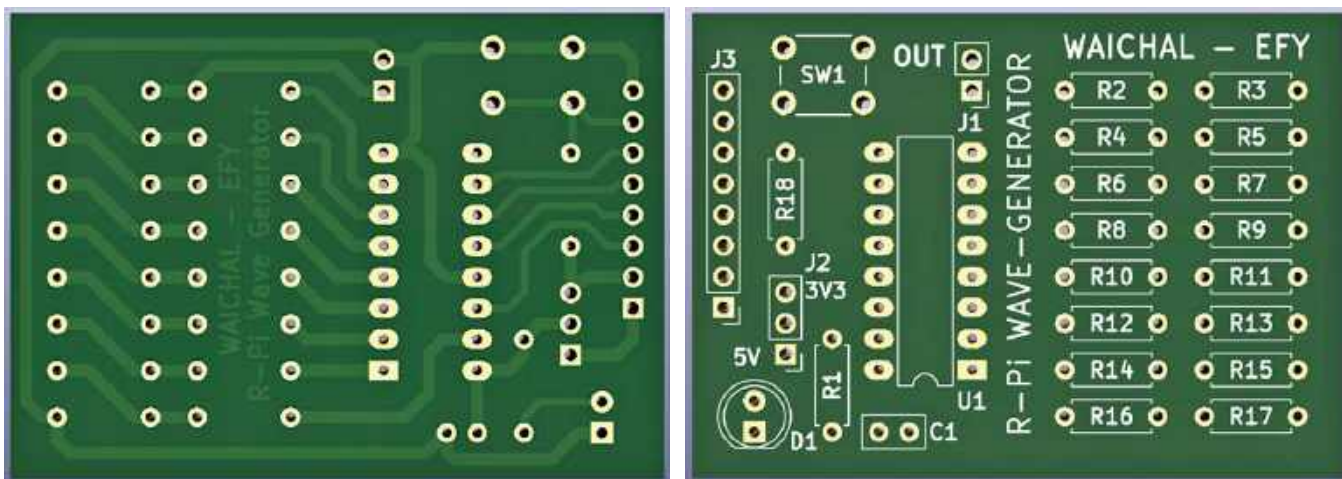
np. do przełączania kształtu fali. To tylko prosty przykład. Użytkownik może dodać więcej klawiszy, napisać własny program i zaimplementować różne funkcje, np. nastawianie częstotliwości i amplitudy.

Uwaga „Electronics For You”: działanie przetwornika DAC zostało przetestowane z programami indywidualnie przygotowanymi dla różnych przebiegów wyjściowych. Stwierdzono przy tym, że zaimplementowanie sterowania klawiaturą zajmuje dużo

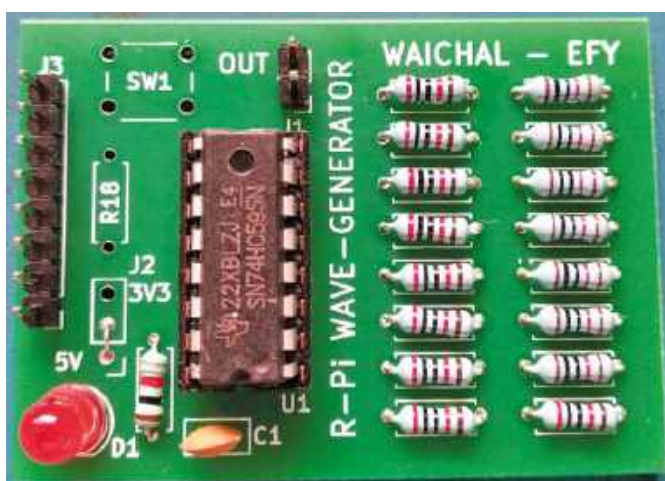
Tabela 2. Połączenia Raspberry Pi z generatorem

generator (J3)	Raspberry Pi	nr styku złącza		Raspberry Pi	generator (J3)
J3/Pin2	+3,3 V	1	2	+5 V	J3/Pin1
	GPIO-2	3	4	+5 V	
	GPIO-3	5	6	GND	J3/Pin8
J3/Pin6	GPIO-4	7	8	GPIO-14	
	GND	9	10	GPIO-15	
	GPIO-17	11	12	GPIO-18	J3/Pin7
	GPIO-27	13	14	GND	
J3/Pin4	GPIO-22	15	16	GPIO-23	
	+3,3 V	17	18	GPIO-24	
J3/Pin3	GPIO-10	19	20	GND	
	GPIO-9	21	22	GPIO-25	
J3/Pin5	GPIO-11	23	24	GPIO-8	
	GND	25	26	GPIO-7	

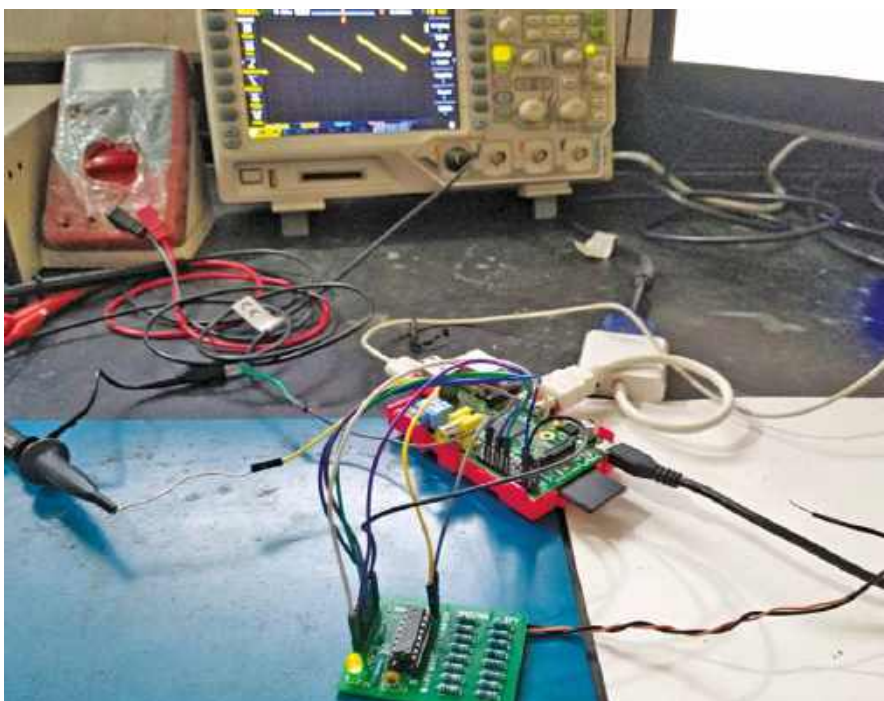
Piny od 27 do 40 nie są używane. Powyższe połączenia są zgodne z wcześniejszymi wersjami płytek Raspberry Pi. Piny używane w generatorze są zaznaczone na kolorowo



Rysunek 3. Płytki drukowane: (a) strona lutowania, (b) strona elementów



Rysunek 4. Zmontowana płytki generatora



Rysunek 5. Testowanie

czasu i może ujemnie wpłynąć na działanie przetwornika cyfrowo-analogowego i na szybkość aktualizacji przebiegu wyjściowego.

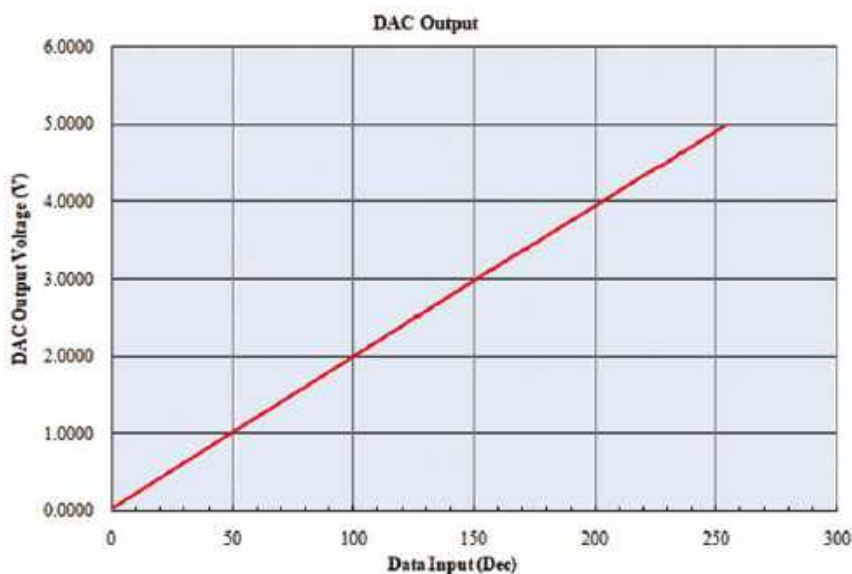
Przypis redaktora: programy dołączone do projektu można uznać za prostą ilustrację podstawowych funkcji generatora. W programach tych przetwornik DAC jest zapisywany w pętli programu głównego, a odmierzenie okresu próbkowania jest realizowane przez wywoływanie podprogramu opóźniającego. Zalecanym sposobem sterowania DAC-a byłoby okresowe wysyłanie do niego danych w podprogramie obsługi przerwań jakiegoś timera. Wówczas w programie głównym dałoby się realizować nawet złożone funkcje bez wpływu na proces generowania sygnału.

Połączenia

Tabela 2 przedstawia połączenia elementów przetwornika DAC generatora ze złączem GPIO płytki Raspberry Pi. Do połączenia obu bloków służy złącze J3. Testy prototypu przeprowadzono z płytką Raspberry Pi w wersji 3, która ma 26-stykowe złącze GPIO. Nowsze wersje, np. Pi 4, mają złącze 40-pinowe. Wersja 4 może być użyta w projekcie, ponieważ piny GPIO w obu wersjach są ze sobą zgodne.

Rysunek 3 przedstawia płytkę drukowaną urządzenia. Pierwszy prototyp został zbudowany i przetestowany na płytce prototypowej. Później zaprojektowano jednostronną płytkę drukowaną i układ na niej również przetestowano. Rysunek 3a pokazuje stronę lutowania płytki, a rysunek 3b – stronę elementów.

Aby uruchomić program w Pythonie, na Raspberry instalujemy Pi Python 3 (jeśli używamy starszej wersji płytki, to Python 2). Jeśli używamy najnowszego systemu operacyjnego Raspberry Pi, to powinien być on fabrycznie zainstalowany na płytce. W materiałach do artykułu znajduje się kilka



Rysunek 6. Wykres liniowości przetwornika DAC

programów w języku Python, przeznaczonych do generowania różnych przebiegów: sinusoidalnego, prostokątnego, piłokształtnego narastającego i opadającego oraz trójkątnego. Każdy program można na Raspberry Pi załadować i uruchomić. Użytkownik może też oczywiście napisać własny program.

Aby program uruchomić, należy źródłowy plik Pythona otworzyć w środowisku IDE

i uruchomić go. Inną możliwością jest otworzyć folder w terminalu i użyć następującego polecenia:

```
sudo python3 nazwa.py
```

Na przykład plik do wytwarzania fali sinusoidalnej sine.py uruchamiamy następująco:

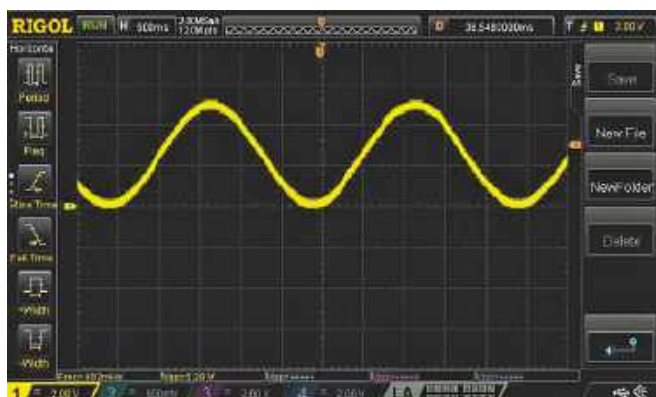
```
sudo python3 sine.py
```



Rysunek 7a. Przebieg wyjściowy/fala piłokształtna narastająca



Rysunek 7b. Przebieg wyjściowy/fala piłokształtna opadająca



Rysunek 7c. Przebieg wyjściowy/fala sinusoidalna



Rysunek 7d. Przebieg wyjściowy/fala prostokątna

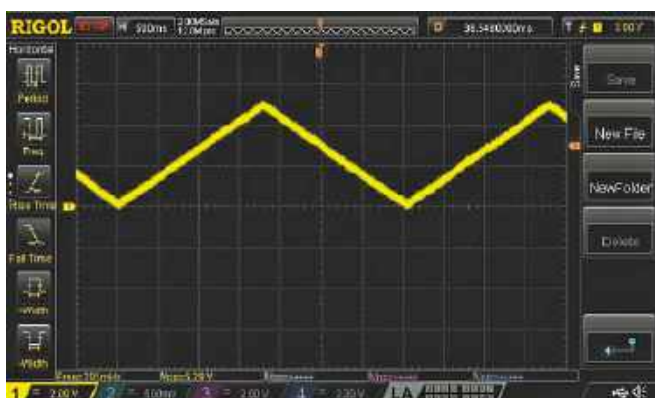
Konstrukcja i testowanie

Montujemy układ na płytce drukowanej.

Rysunek 4 przedstawia zmontowaną płytkę profesjonalnej jakości. Można na początek zmontować tylko układ przetwornika cyfrowo-analogowego R-2R i przetestować jego zachowanie dla różnych słów wejściowych w zakresie od 0 do 255 (00000000₂ do 11111111₂). Jest to proces czasochłonny, ale pozwoli sprawdzić przetwornik pod kątem jego liniowości, co na późniejszym etapie zagwarantuje prawidłowe generowanie przebiegów. W przypadku jakichkolwiek problemów z przetwornikiem, testowanie szybko je ujawni i pozwoli zaoszczędzić czas w dalszym etapie uruchamiania.

Rysunek 5 przedstawia układ podczas testów, natomiast **rysunek 6** ilustruje liniowość przetwornika cyfrowo-analogowego dla słów wejściowych z zakresu 0...255.

Po sprawdzeniu działania przetwornika DAC można zmontować i przetestować pozostałą część generatora. Przed podłączeniem Raspberry Pi należy się upewnić, że na płytce nie ma zwarc. Podłączamy zasilanie +3,3 V, +5 V i GND z płytki Pi do płytki drukowanej generatora, sprawdzamy, czy nie ma zwarc na liniach zasilania, a następnie realizujemy pozostałe połączenia obu płytek. Pierwsze załączenie zasilania



Rysunek 7e. Przebieg wyjściowy/fala trójkątna



Rysunek 7f. Komunikacja szeregową po SPI: kanał żółty – SCLK, kanał niebieski – SDATA, kanał fioletowy – strobe/latch

przeprowadzamy bez układu scalonego U1. Sprawdzamy, czy dioda LED D1 świeci prawidłowo. Mierzmy napięcie doprowadzone do pinów 16 (+5 V) i 8 (GND) układu U1. Gdy są one prawidłowe, wyłączamy zasilanie i montujemy układ U1, po czym znów załączamy zasilanie.

W środowisku programowania Python na Raspberry Pi ładujemy program przeznaczony do generowania fali trójkątnej i uruchamiamy go. Obserwujemy na oscyloskopie przebieg wyjściowy na złączu J1 (DAC_OUT). Jeśli przebieg jest prawidłowy,

można przetestować programy wytwarzające przebiegi sinusoidalny, prostokątny i piłokształtny, a także uruchamiać własne programy, generujące dowolne inne przebiegi.

Podczas testów zaobserwowano, że tempo obsługi DAC-a jest niezbyt szybkie, co sprawia, że generator nadaje się raczej do wytwarzania przebiegów niskiej częstotliwości, nawet w zakresie pojedynczych herców. Ograniczenie to może wynikać z działania istniejących bibliotek i funkcji obsługiwanych w środowisku Pythona na Raspberry Pi.

Rysunek 7 ukazuje różne przebiegi (A...F) wygenerowane przy użyciu generatora i przedstawia oscylogramy odbierania danych z portu SPI przez rejestr przesuwający. ■

Prasanna Waichal, Arpita Waichal, Nikita Thorat, Pragati Todkar

• Materiał filmowy do artykułu:
<https://youtu.be/wt1dVr76rls>

Materiały dodatkowe są dostępne na stronie elportal.pl/do-pobrania

Ten artykuł był wcześniej opublikowany na łamach „EFY”, październik 2024 (efymag.com)

REKLAMA

numery archiwalne • prenumerata • książki
www.UlubionyKiosk.pl



Tymek – Młodzi Entuzjaści Elektroniki, Wrocław

Serce to jeden z najbardziej rozpoznawalnych symboli. Kojarzy się z życiem, emocjami i rytmem, który towarzyszy nam każdego dnia. Przede wszystkim kojarzy się jednak z miłością. W walentynkowym czasie szczególnie często pojawia się na kartkach, prezentach i dekoracjach, zwykle w towarzystwie migoczących światełek. W elektronice również potrafi „bić”, choć robi to dokładnie tak, jak zostało zaprogramowane – regularnie, powtarzalnie i bez niespodzianek. I właśnie takie elektroniczne serce zbudujemy podczas tego spotkania.

Miłość ma wiele twarzy i nie zawsze wygląda tak samo. Czasem jest cicha i spokojna, a czasem radosna i pełna energii. Bywa widoczna od razu, ale zdarza się też, że kryje się w drobnych gestach, które łatwo przeoczyć. Właśnie dlatego tak trudno ją jednoznacznie opisać, a jednocześnie tak łatwo ją poczuć.

Najczęściej myślimy o miłości w kontekście jednej, bardzo bliskiej osoby. Kogoś, przy kim czujemy się bezpiecznie i swobodnie. Kogoś, z kim możemy rozmawiać, śmiać się i dzielić codzienne sprawy – nawet te zupełnie zwyczajne. Taka miłość potrafi dodawać odwagi i sprawiać, że nawet te trudne dni stają się łatwiejsze.

Jest jednak także miłość, która towarzyszy nam od najmłodszych lat. To miłość

rodziców i opiekunów – cierpliwa, wyrozumiała i często niewidoczna na pierwszy rzut oka. Objawia się w trosce, w pomocy, w uważnym towarzyszeniu i w zasadach, które chronią, a nie ranią, oraz w obecności wtedy, gdy najbardziej jej potrzebujemy.

Jest też miłość do dzieci. Pełna dumy, radości i chęci chronienia ich przed tym, co trudne i niebezpieczne. To miłość, która uczy odpowiedzialności i pokazuje, jak ważne są ciepło, uwaga i czas spędzony razem.

Nie można zapominać o miłości do samego siebie. To ona pozwala nam zaakceptować swoje błędy i niedoskonałości. Uczy, że warto dbać o własne potrzeby, słuchać siebie i nie pozwalać innym przekraczać naszych granic. Bez niej trudno budować zdrowe relacje z innymi.

Miłość potrafi wyjść poza dom i rodzinę. Widać ją w pomaganiu innym, w wolontariacie, w drobnych gestach życzliwości wobec osób, których nawet dobrze nie znamy. To ona sprawia, że chcemy zrobić coś dobrego dla świata – choćby małymi krokami.

Bywa także miłością do przyrody, do nauki, do pasji i zainteresowań. To dzięki niej chcemy poznawać, odkrywać i tworzyć. To ona często prowadzi nas do osiągnięć i rzeczy, które wcześniej istniały tylko w wyobraźni.

Czasem mówi się nawet o miłości do wrogów. Nie oznacza to zgody na zło ani zapomniania o krzywdach, lecz próbę nieulegania nienawiści, która potrafi odebrać spokój i radość. To trudna, ale bardzo dojrzała postawa.



Fotografia 1. Od lewej: Marcel i Tymek podczas montażu zestawu Bijące serce LED (AVTEDU620). Młodzi Entuzjaści Elektroniki, Zagroda CUDów, Trójca

Miłość w najlepszym wydaniu dodaje skrzydeł i sprawia, że świat wydaje się jaśniejszy i bardziej przyjazny. Warto jednak pamiętać, że nie każde uczucie, które nosi taką nazwę, naprawdę nią jest, i że nawet najpiękniejsza miłość wymaga uważności, rozsądku oraz troski o samego siebie.

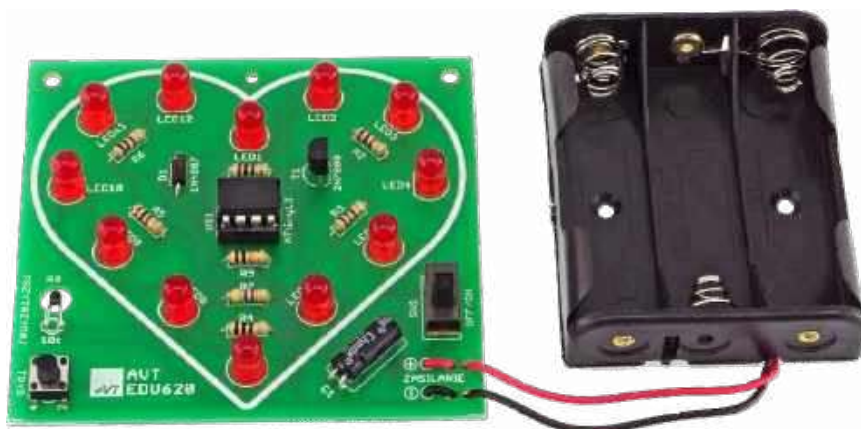
Co zbudujemy tym razem?

Podczas tego spotkania zbudujemy zestaw **AVTEDU620 – Bijące serce LED** – niewielki układ, który potrafi ożywać światłem. Diody LED nie migają tu przypadkowo, lecz rozjaśniają się i gasną w określonym rytmie, tworząc wrażenie pulsowania przypominającego bicie serca. Charakterystyczną cechą tego projektu jest to, że tempo ustawionego efektu (do wyboru są trzy) nie jest stałe. Układ reaguje na temperaturę czujnika – gdy dotkniemy go palcem lub zbliżymy do źródła ciepła, rytm światła zmienia się i staje się szybszy. Dzięki temu serce LED sprawia wrażenie, jakby reagowało na otoczenie. Taki sposób działania sprawia, że gotowy układ jest nie tylko efektowny, ale też ciekawy w obserwacji. Można sprawdzić, jak zmienia się jego „puls”, porównać różne sytuacje i przekonać się, jak subtelnie zastosowanie prostego czujnika pozwala elektronicznie reagować na nasz dotyk.

Na **fotografii 2** pokazano zmontowany układ, a na krótkim filmie pod adresem <https://youtu.be/kLbRc91aqjo> możesz zobaczyć, co dzieje się, gdy elektroniczne serce „czuje” ciepło dłoni.

Schemat montażowy

Schemat montażowy to rysunek, który pokazuje, gdzie dokładnie na płytce drukowanej należy umieścić każdy z elementów zestawu. Dzięki niemu łatwo odnaleźć właściwe miejsce dla rezystorów, kondensatorów, diod czy układów scalonych i innych podzespołów, ponieważ wszystkie komponenty są oznaczone takimi samymi desygntorami, zarówno na schemacie montażowym, liście elementów jak i na schemacie ideowym. Ułatwia to bezbłędne i szybkie składanie układu, nawet osobom początkującym. Schemat montażowy pomaga również uniknąć pomyłek, takich jak wlutowanie elementu w niewłaściwe miejsce lub ustawienie go w złej orientacji. Schemat montażowy, który dodatkowo pokazuje układ ścieżek i padów, bardzo pomaga w kontroli poprawności wykonanych połączeń lutowanych. Dzięki temu łatwo ustalić, czy połączenia pomiędzy sąsiednimi polami są przewidziane w projekcie, czy też powstały przez pomyłkę, na przykład na skutek



Fotografia 2. Bijące serce LED (kod AVTEDU620). Zmontowany układ

przypadkowego zwarcia ich cyną podczas nieostrożnego lutowania. Taki podgląd znacząco ułatwia wykrywanie błędów i zwiększa pewność, że układ został zmontowany prawidłowo. Schemat montażowy pokazano na **rysunku 1**.

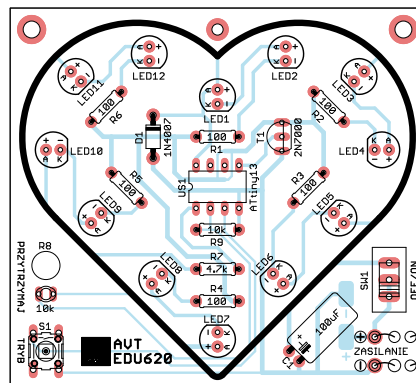
Montaż rezystorów

Zgodnie z informacjami z listy elementów przylutuj rezystory o określonych wartościach rezystancji na odpowiednich pozycjach na płytce drukowanej. Przed przystąpieniem do tej czynności zapoznaj się ze wskazówkami poniżej.

- Rezystor jest elementem, który ogranicza przepływ prądu w obwodzie elektrycznym. Dzieje się tak dlatego, że ma on określoną rezystancję – właściwość materiału wywołującą spadek napięcia podczas przepływu prądu. Energia elektryczna zamienia się w nim na ciepło, co jest naturalnym skutkiem przepływu prądu przez rezystancję.
- Rezystor nie ma biegunowości – działa tak samo niezależnie od kierunku przepływu prądu. Dlatego jego montaż na płytce nie wymaga uwzględnienia orientacji. Ważne jest jedynie, aby w danym miejscu umieścić właściwy rezystor o odpowiedniej wartości. Sam kierunek montażu pozostaje dowolny.
- Rezystory są zazwyczaj jednymi z najniższych elementów montowanych na płytce drukowanej. Ich przylutowanie nie utrudnia późniejszego montażu wyższych komponentów, dlatego lutuje się je zazwyczaj w pierwszej kolejności.
- Spośród wszystkich komponentów znajdujących się w zestawie wyodrębnił rezystory i odłóż je na osobną stertę.
- Pozostałe elementy odłóż na bok, a podczas montażu sięgaj wyłącznie po kolejne rezystory z wcześniej przygotowanej sterty.

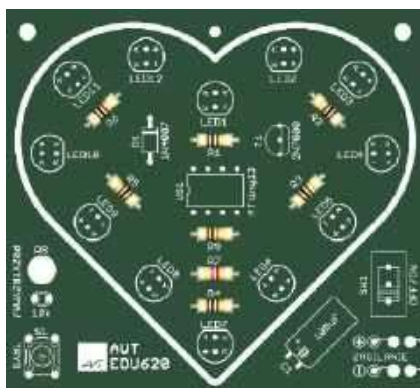
Wykaz elementów:

R1...R6: 100 Ω (brązowy-czarny-brązowy-żółty)
R7: 4,7 kΩ (żółty-fioletowy-czerwony-żółty)
R9: 10 kΩ (brązowy-czarny-pomarańczowy-żółty)
D1: 1N4007
T1: 2N7000
C1: 100 μF
LED1...LED12: dioda LED czerwona
S1: mikroprzycisk
U1: ATtiny13 + podstawka
SW1: włącznik
R8: termistor 10 kΩ
złączka do baterii: czerwony +, czarny -



Rysunek 1. Schemat montażowy układu

- Za każdym razem, gdy weźmiesz do ręki kolejny rezystor, zmierz jego wartość za pomocą multimetru ustawionego w tryb omomierza. Odczytaną rezystancję zapamiętaj.
- Jeśli pomiar rezystorów sprawia Ci trudność, poproś o pomoc kolegę lub osobę prowadzącą zajęcia. Gdy masz dostęp do internetu, możesz też skorzystać z instrukcji dostępnej na stronie <https://elportal.pl/do-pobrania> – znajdziesz tam dokument „Pomiar wartości rezystorów za pomocą multimetru”, przygotowany jako materiał uzupełniający do EdW 11/2024.
- Na dołączonej do zestawu liście elementów odszukaj zmierzoną wcześniej wartość rezystancji (w Ω, kΩ lub MΩ),



Rysunek 2. Rezystory o odpowiednich wartościach, oznaczone za pomocą kodu kolorowych pasków, zamontowane na właściwych pozycjach, zgodnie z listą elementów

a następnie odczytaj desygnator lub desygnatory przypisane do tej wartości.

- Zegnij wyprowadzenia rezystora i umieść go w płytce (patrz rysunek) w miejscu oznaczonym właściwym desygnatorem.
- Zadbaj o to, aby każdy rezystor był włożony do płytki do końca i dobrze do niej przylegał. Estetyczny montaż nie tylko poprawia wygląd gotowego urządzenia, lecz także stabilizuje element w płytce, chroniąc go przed uszkodzeniami mechanicznymi, oraz ułatwia późniejszą diagnostykę i ewentualne naprawy.
- Przyłutuj element do płytki. Jeśli nie wiesz, jak to zrobić, albo chcesz upewnić się, że wykonujesz to prawidłowo, przeczytaj sekcję *Lutowanie komponentów przewlekanych do płytki drukowanej*. Nieco poniżej znajdziesz również informacje, jak wykonać tę czynność w sposób bezpieczny dla siebie i pozostałych uczestników zajęć.
- Usuń nadmiar wyprowadzeń za pomocą obcinaczek. Jeśli nie wiesz, jak to zrobić, albo chcesz upewnić się, że wykonujesz to prawidłowo, przeczytaj sekcję *Przycinanie nadmiaru wyprowadzeń*. Nieco poniżej znajdziesz również informacje, jak wykonać tę czynność w sposób bezpieczny dla siebie i pozostałych uczestników zajęć.

Montaż podstawki pod mikrokontroler

Przylutuj do płytki podstawkę pod układ US1. Przed przystąpieniem do tej czynności zapoznaj się ze wskazówkami poniżej.

- Podstawka pod układ scalony (tu mikrokontroler) przewodzi prąd pomiędzy wyprowadzeniami układu a ścieżkami na płytce, ale sama w sobie nie pełni żadnej aktywnej funkcji elektrycznej – nie zmienia sygnałów, nie wzmacnia ich ani nie ogranicza. Jej głównym zadaniem jest zapewnienie wygodnego, bezpiecznego i wielokrotnego montażu układu scalonego bez ryzyka uszkodzenia jego wyprowadzeń lub pól lutowniczych na płytce. Podstawka umożliwia między innymi łatwą wymianę układu scalonego w razie jego uszkodzenia lub pomyłki podczas montażu.
- Podstawka pod układ scalony nie ma biegunowości w sensie elektrycznym, ale ma określoną orientację montażową. W jej obudowie znajduje się znacznik (najczęściej wycięcie lub kropka), który musi być ustawiony zgodnie ze znakiem na płytce drukowanej (**rysunek 3**). Prawidłowa orientacja podstawki jest konieczna, aby później poprawnie włożyć do niej układ scalony.
- Podstawki pod układy scalone należy do elementów o średniej wysokości, dlatego zazwyczaj montuje się je po przylutowaniu najniższych komponentów, takich jak rezystory. Dzięki temu płytka pozostaje stabilna, a podstawka nie zasłania dostępu do miejsc montażowych przeznaczonych dla pozostałych elementów.
- Przyłutuj element do płytki. Jeśli nie wiesz, jak to zrobić, albo chcesz upewnić się, że wykonujesz to prawidłowo, przeczytaj sekcję *Lutowanie komponentów przewlekanych do płytki drukowanej*. Nieco poniżej znajdziesz również informacje, jak wykonać tę czynność w sposób bezpieczny dla siebie i pozostałych uczestników zajęć.
- Podczas lutowania pinów do płytki staraj się, by pomiędzy lutowanymi

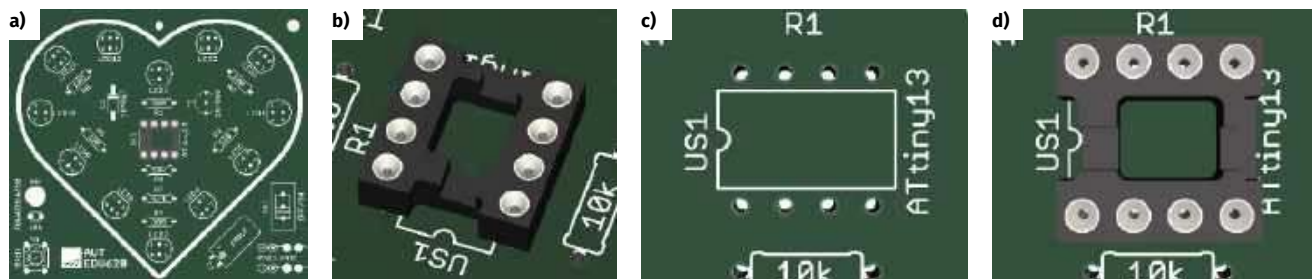
wyprowadzeniami nie powstały niechciane połączenia, czyli zwarcia, zwane również mostkami lutowniczymi. Jeśli podczas lutowania pojawiają się zwarcia, najłatwiej będzie, trzymając płytkę jedną ręką, ustawić ją pod kątem prostym względem blatu. Następnie należy ponownie podgrzać połączone pola lutownicze oraz przy pomocy grotu lutownicy i siły grawitacji pozwolić nadmiarowi cyny spłynąć na blat. Dzięki temu uwolnisz pady podstawki od zwarć.

- W przypadku podstawek nie ma potrzeby przycinania wyprowadzeń. Po przylutowaniu pozostaw je w oryginalnej długości.

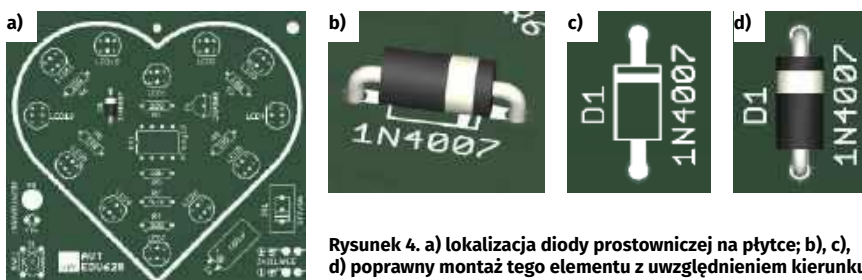
Montaż diody prostowniczej

Kolejnym elementem do zamontowania jest dioda prostownicza 1N4007. Należy zastosować sposób postępowania jak dla diod D1 i D2, opisany powyżej.

- Dioda prostownicza jest elementem, który przewodzi prąd głównie w jednym kierunku, a blokuje jego przepływ w kierunku przeciwnym. Wynika to z właściwości złącza półprzewodnikowego, które przy odpowiedniej polaryzacji przewodzi, a przy odwrotnej – pozostaje zamknięte. Dzięki temu diody umożliwiają sterowanie kierunkiem przepływu prądu, ochronę układów oraz przetwarzanie sygnałów elektrycznych.
- Dioda ma biegunowość – przewodzi prąd tylko w jednym kierunku. Dlatego jej montaż na płytce wymaga zachowania właściwej orientacji. Na obudowie diody znajduje się pasek oznaczający katodę, który musi być ustawiony zgodnie ze znakiem na płytce drukowanej. Umieszczenie diody w złym kierunku spowoduje, że nie będzie działała prawidłowo w układzie.
- Diody są z reguły elementami o niewielkiej wysokości, dlatego również montuje się je we wczesnym etapie lutowania – zwykle tuż po rezystorach



Rysunek 3. a) lokalizacja podstawki na płytce; b), c), d) poprawny montaż tego elementu z uwzględnieniem kierunku



Rysunek 4. a) lokalizacja diody prostowniczej na płytce; b), c), d) poprawny montaż tego elementu z uwzględnieniem kierunku

lub podstawkach. Dzięki temu płytka pozostaje stabilna, a diody nie utrudniają późniejszego montażu wyższych komponentów.

- Za każdym razem, gdy weźmiesz do ręki kolejną diodę, odczytaj jej oznaczenie na obudowie i porównaj je z informacją w liście elementów. Diody różnią się między sobą typami i parametrami, dlatego ważne jest, aby upewnić się, że właściwa dioda zostanie zamontowana we właściwym miejscu na płytce.
- Warto zadbać o to, aby każda dioda była włożona do płytki do końca i dobrze do niej przylegała. Estetyczny montaż nie tylko poprawia wygląd gotowego urządzenia, lecz także stabilizuje element w płytce, chroniąc go przed uszkodzeniami mechanicznymi, oraz ułatwia późniejszą diagnostykę i ewentualne naprawy.
- Jeśli rozpoznawanie typów diod sprawia Ci trudność, poproś o pomoc kolegę lub osobę prowadzącą zajęcia. Możesz też spróbować poszukać informacji na ten temat w Internecie, w książkach lub u agentów AI – pod warunkiem, że masz taką możliwość i potrafisz zweryfikować wiarygodność informacji z tych źródeł.
- Zegnij wyprowadzenia diody i umieść ją w płytce (rysunek 4) w miejscu oznaczonym właściwym desygnatorem, pamiętając o zachowaniu jej orientacji (patrz wyżej).
- Przylutuj element do płytki. Jeśli nie wiesz, jak to zrobić, albo chcesz upewnić się, że wykonujesz to prawidłowo, przeczytaj sekcję *Lutowanie*

komponentów przewlekanych do płytki drukowanej. Nieco poniżej znajdziesz również informacje, jak wykonać tę czynność w sposób bezpieczny dla siebie i pozostałych uczestników zajęć.

- Usuń nadmiar wyprowadzeń za pomocą obcinaczek. Jeśli nie wiesz, jak to zrobić, albo chcesz upewnić się, że wykonujesz to prawidłowo, przeczytaj sekcję *Przycinanie nadmiaru wyprowadzeń.* Nieco poniżej znajdziesz również informacje, jak wykonać tę czynność w sposób bezpieczny dla siebie i pozostałych uczestników zajęć.

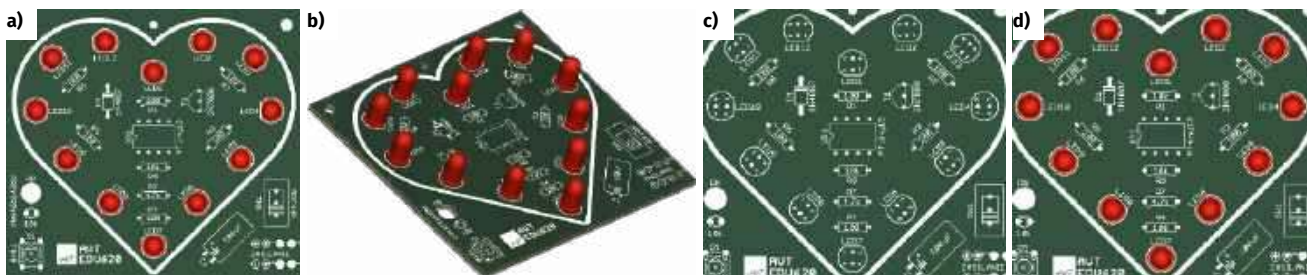
Montaż diod LED

Zgodnie z informacjami z listy elementów na odpowiednich pozycjach przylutuj czerwone diody LED. Przed przystąpieniem do tej czynności zapoznaj się ze wskazówkami poniżej.

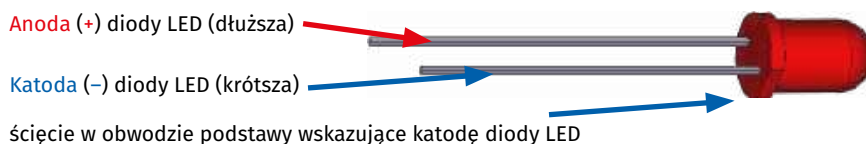
- Dioda LED to element elektroniczny, który świeci, gdy płynie przez niego prąd w odpowiednim kierunku. Łączy w sobie działanie zwykłej diody – przewodzi prąd tylko w jedną stronę – oraz funkcję źródła światła. Dzięki temu LED-y mogą sygnalizować działanie układu, informować o stanie pracy urządzenia lub pracować w układach generujących efekty świetlne.
- Tak jak każda dioda, LED ma biegunowość. Oznacza to, że musi być podłączona we właściwym kierunku, inaczej nie zaświeci, a w szczególnym przypadku ulegnie uszkodzeniu. Jej katodę najczęściej oznacza ścięcie na obudowie oraz krótsza nóżka (rysunek 6). Przed montażem sprawdź, gdzie na PCB

znajduje się oznaczenie katody, i ustaw diodę zgodnie z nim.

- Dioda LED jest elementem, który od razu przyciąga wzrok obserwatora, dlatego estetyka jej montażu ma duży wpływ na końcowy wygląd budowanego urządzenia. Warto zadbać o to, aby LED była ustawiona prostopadle do płytki i równo do niej przylegała – nawet drobne odchylenia mogą być widoczne po uruchomieniu układu, szczególnie gdy diod jest więcej.
- Z uwagi na powyższe LED-y najlepiej montować na stosunkowo wczesnym etapie lutowania. Rezystory, podstawki pod układy scalone oraz diody prostownicze i sygnałowe są zazwyczaj nieco niższe, ale zaraz po nich warto umieścić na płytce diody LED. W tym momencie pole lutownicze jest wciąż dobrze dostępne, i nic nie zasłania miejsca montażu, co ułatwia przylutowanie LED-ów równo i estetycznie.
- Zanim włożysz diodę LED do płytki, sprawdź w liście elementów, jaki kolor LED-a powinien zostać zamontowany w danej lokalizacji. Same diody – zwłaszcza w bezbarwnych obudowach – mogą wyglądać bardzo podobnie lub wręcz identycznie, dlatego warto upewnić się, jaki kolor świecenia ma LED, który trzymasz w ręku.
- Jeśli w projekcie występuje kilka kolorów diod LED w bezbarwnych obudowach, zasadne jest ich wcześniejsze posegregowanie. Najprościej zrobić to za pomocą multimetru ustawionego w tryb badania diod lub ciągłości obwodu. Przyłożenie sond – czerwonej do anody diody LED i czarnej do jej katody (fotografia 3) – spowoduje lekkie świecenie diody LED, co pozwoli od razu ustalić jej kolor. Dzięki temu można przyporządkować poszczególne LED-y do właściwych grup i ułożyć je na osobnych stertach. Takie przygotowanie znacząco zmniejsza ryzyko pomyłek podczas montażu i gwarantuje



Rysunek 5. a) lokalizacja diod LED na płytce; b), c), d) poprawny montaż elementów z uwzględnieniem kierunku



Rysunek 6. Opis wyprowadzeń diody LED („plusowe” wyprowadzenie dłuższe, „minusowe” krótsze)



Fotografia 3. Sprawdzanie diody LED za pomocą multimetru ustawionego na funkcję testowania diod. Po przyłożeniu sondy czerwonej do anody, a czarnej do katody, sprawna dioda LED powinna się zaświecić. Jeśli dioda ma odpowiednio długie (jeszcze nie przycięte) wyprowadzenia można się wspomóc krokodylkami

prawidłowy efekt wizualny w gotowym urządzeniu.

- Jeśli upewniłeś się co do odpowiedniej polaryzacji i kolorów, przylutuj wcześniej włożone diody LED (lub każdą z osobna) do płytki drukowanej. Jeśli nie wiesz, jak to zrobić, albo chcesz upewnić się, że wykonujesz to prawidłowo, przeczytaj sekcję *Lutowanie komponentów przewlekanych do płytki drukowanej*. Nieco poniżej znajdziesz również informacje, jak wykonać tę czynność w sposób bezpieczny dla siebie i pozostałych uczestników zajęć.
- Usuń nadmiar wyprowadzeń za pomocą obcinaczek. Jeśli nie wiesz, jak to zrobić, albo chcesz upewnić się, że wykonujesz to prawidłowo, przeczytaj sekcję *Przycinanie nadmiaru wyprowadzeń*. Nieco poniżej znajdziesz również informacje, jak wykonać tę czynność w sposób bezpieczny dla siebie i pozostałych uczestników zajęć.

Montaż tranzystora polowego N-MOSFET

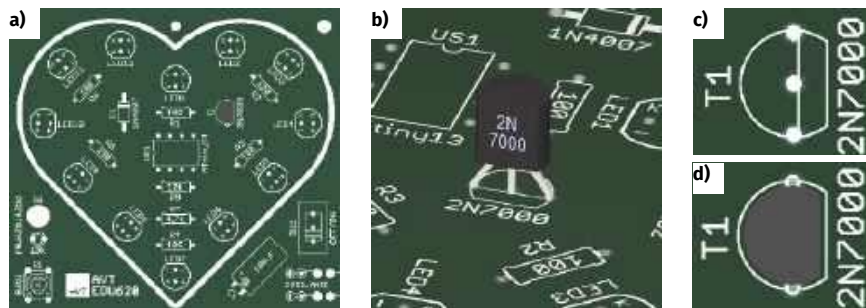
Zgodnie z informacjami z listy elementów przylutuj tranzystor N-MOSFET, T1 typu 2N7000 we wskazanej lokalizacji. Przed przystąpieniem do tej czynności zapoznaj się ze wskazówkami poniżej.

- Tranzystor polowy (MOSFET) to element półprzewodnikowy umożliwiający sterowanie przepływem prądu pomiędzy drenem a źródłem poprzez przyłożenie napięcia do bramki. Może pracować jako wzmacniacz sygnału lub jako klucz elektroniczny sterujący innymi elementami, takimi jak diody LED, przekaźniki czy brzęczki.

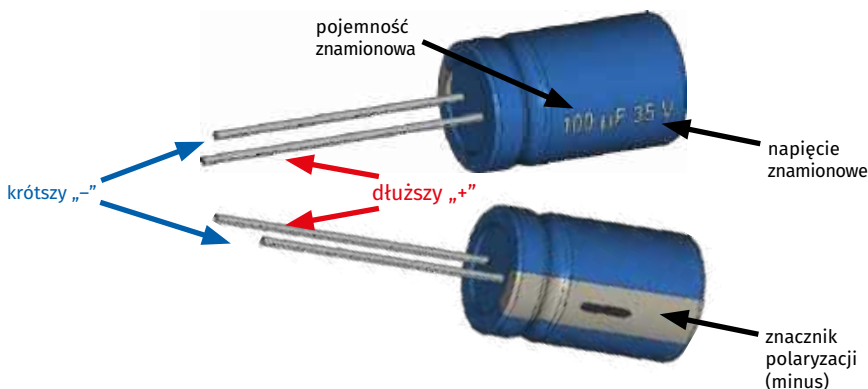
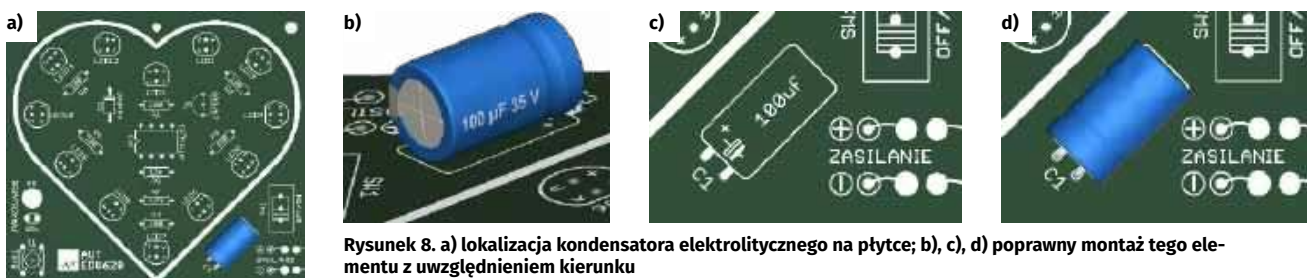
- W praktyce stosuje się tranzystory N-MOSFET i P-MOSFET, które różnią się polaryzacją napięć sterujących oraz kierunkiem przewodzenia prądu w układzie. Konkretny typ tranzystora jest zawsze określony w liście elementów i na schemacie.
- W układach elektronicznych stosuje się tranzystory polowe o określonych parametrach i przeznaczeniu. Przed montażem należy sprawdzić w liście elementów, jaki tranzystor jest przewidziany w danej lokalizacji, i zastosować dokładnie ten typ lub zalecany zamiennik.
- Tranzystor polowy ma trzy wyprowadzenia: bramkę (G), dren (D) i źródło (S). Ich rozmieszczenie zależy od typu obudowy i producenta, dlatego przed montażem należy porównać układ nóżek z oznaczeniem na płytce drukowanej oraz z dokumentacją elementu. Prawidłowa orientacja tranzystora jest warunkiem jego poprawnej pracy.
- Bramka tranzystora MOSFET jest wrażliwa na wyładowania elektrostatyczne (ESD). Podczas montażu należy unikać dotykania wyprowadzeń palcami oraz pracować w możliwie bezpiecznych warunkach antystatycznych. Zaleca się lutowanie tranzystora jako jednego z ostatnich elementów.
- Tranzystory polowe (podobnie jak bipolarnie) mają charakterystyczny kształt obudowy (najczęściej półokrągły z jednej strony), co ułatwia ich prawidłowe ustawienie. Element należy włożyć do płytki tak, aby płaska i zaokrąglona część obudowy odpowiadała oznaczeniu na PCB. Obrys

komponentu na płytce pomaga w szybkim i poprawnym montażu.

- Tranzystory polowe montuj po elementach najniższych, takich jak rezystory i diody, a przed komponentami wysokimi, na przykład kondensatorami elektrolitycznymi czy złączami. Zapewnia to wygodny dostęp do wyprowadzeń podczas lutowania oraz stabilne ustawienie elementu.
- Zadbaj o estetyczny montaż tranzystora. Obudowa powinna być ustawiona pionowo, a wyprowadzenia równomiernie przylegać do płytki drukowanej. Poprawne ustawienie ułatwia lutowanie i zwiększa odporność na uszkodzenia mechaniczne.
- Po umieszczeniu tranzystora w płytce delikatnie odegnij dwie skrajne nóżki, aby pozostał na miejscu podczas lutowania. Najpierw przylutuj środkowe wyprowadzenie, następnie ustaw element do pionu i przylutuj pozostałe. Lutuj krótko i sprawnie, aby niepotrzebnie nie nagrzewać wyprowadzeń.
- Jeśli chcesz upewnić się, że lutowanie wykonujesz prawidłowo, zapoznaj się z sekcją *Lutowanie komponentów przewlekanych do płytki drukowanej*. Nieco poniżej znajdziesz również informacje, jak wykonać tę czynność w sposób bezpieczny dla siebie i pozostałych uczestników zajęć.
- Po zakończeniu lutowania usuń nadmiar wyprowadzeń za pomocą obcinaczek. Wskazówki dotyczące tej czynności znajdziesz w sekcji *Przycinanie nadmiaru wyprowadzeń*, wraz z zasadami bezpiecznej pracy.



Rysunek 7. a) lokalizacja tranzystora na płytce; b), c), d) poprawny montaż tego elementu z uwzględnieniem kierunku



Rysunek 9. Na korpusie kondensatora elektrolitycznego odnajdziesz – między innymi – informacje o nominalnej pojemności oraz dopuszczalnym napięciu pracy a także znacznik polaryzacji

Montaż kondensatora elektrolitycznego

Zgodnie z informacjami z listy elementów przylutuj kondensator elektrolityczny C1 na wskazanej lokalizacji. Przed przystąpieniem do tej czynności zapoznaj się ze wskazówkami poniżej.

- Kondensator elektrolityczny to element spolaryzowany, który może magazynować stosunkowo duży ładunek elektryczny i pełnić w układzie różne funkcje: stabilizować napięcie, wygładzać tętnienia, filtrować zakłócenia lub dostarczać krótkotrwałych impulsów prądowych. Dzięki dużej pojemności w niewielkiej obudowie jest często stosowany w zasilaczach i układach energoelektronicznych.
- Kondensatory elektrolityczne mają zawsze określoną biegunowość. Na obudowie znajduje się wyraźne oznaczenie minusa (zwykle biały pasek), a dłuższa noga oznacza plus zasilania. Montaż odwrotny

grozi uszkodzeniem kondensatora, a w skrajnych przypadkach nawet jego rozerwaniem (rozszczerzeniem i dezintegracją). Zawsze upewnij się, że plus i minus znajdują się we właściwych otworach na płytce.

- Przed montażem sprawdź zgodność pojemności i napięcia kondensatora z miejscem, w którym ma zostać umieszczony. Odczytaj nadruk z obudowy (na przykład „100 μ F 35 V”) i porównaj go z informacją w liście elementów.
- Kondensatory elektrolityczne są wysokimi elementami, dlatego montuje się je dopiero po wlutowaniu wszystkich niższych komponentów, takich jak rezystory, diody, tranzystory czy kondensatory foliowe. Taka kolejność ułatwia pracę oraz zapewnia stabilne oparcie płytki podczas lutowania.
- Po umieszczeniu kondensatora w otworach sprawdź jeszcze raz jego orientację. W przypadku elementów spolaryzowanych warto wyrobić sobie

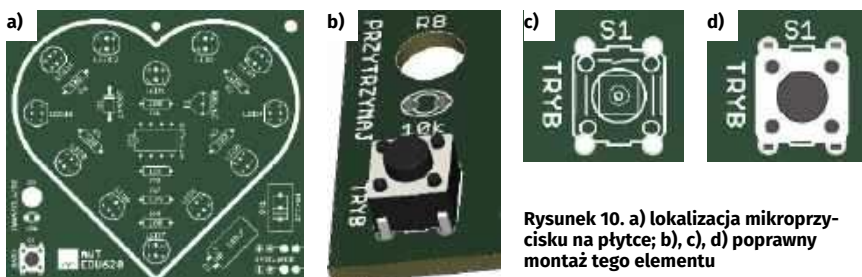
nawyk podwójnego sprawdzenia przed przylutowaniem – to pozwala uniknąć potencjalnie niebezpiecznych błędów.

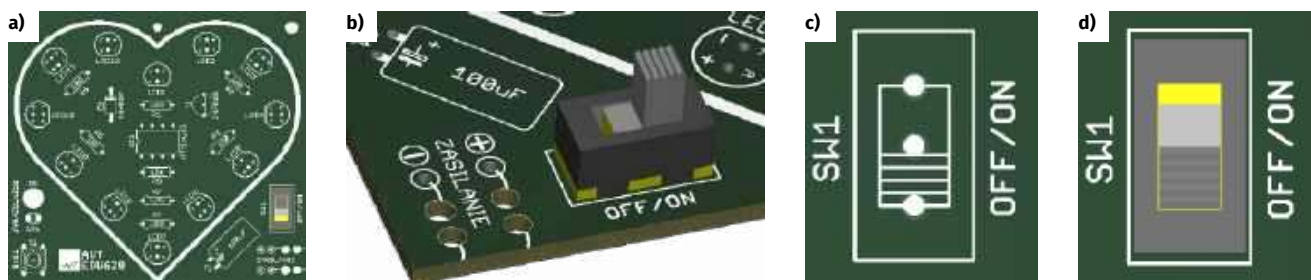
- Aby kondensator nie wypadł podczas obracania płytki, delikatnie odchyl jego wyprowadzenia na zewnątrz. Wystarczy niewielkie odgięcie pod kątem kilku stopni. Zbyt silne zaginanie może utrudnić przycinanie nadmiaru wyprowadzeń, a także przyszły ewentualny demontaż kondensatora.
- Przylutuj kondensator elektrolityczny do płytki, dbając o solidne, czyste punkty lutownicze. W razie wątpliwości zajrzyj do sekcji *Lutowanie komponentów przewlekanych do płytki drukowanej*, gdzie opisano właściwą technikę lutowania i zasady bezpieczeństwa.
- Po zakończeniu lutowania odetnij nadmiar wyprowadzeń obcinaczkami. Wskazówki dotyczące bezpiecznego przycinania znajdziesz w sekcji *Przycinanie nadmiaru wyprowadzeń*.
- Pamiętaj o bezpieczeństwie. Kondensator elektrolityczny zamontowany odwrotnie lub podłączony do wyższego niż znamionowe napięcie może ulec uszkodzeniu, a nawet gwałtownie wybuchnąć. Dlatego przed pierwszym podłączeniem zasilania zawsze upewnij się, że został zamontowany poprawnie i ma właściwe parametry.
- Przed pierwszym podłączeniem napięcia do układu zawierającego kondensatory elektrolityczne obowiązkowo zakładaj okulary ochronne.

Montaż mikroprzycisku (tact switch)

Zgodnie z informacjami z listy elementów zamontuj mikroprzełącznik S1 na wskazanej lokalizacji. Przed przystąpieniem do tej czynności zapoznaj się ze wskazówkami poniżej.

- Mikroprzycisk (tact switch) to chwilowy przełącznik, który zwiera styki tylko w czasie naciśnięcia przycisku. Po zwolnieniu nacisku automatycznie wraca do stanu początkowego. Najczęściej wykorzystywany jest jako





Rysunek 11. a) lokalizacja przełącznika na płytce; b), c), d) poprawny montaż tego elementu

przycisk sterujący, reset lub przycisk funkcyjny.

- Typowy mikroprzycisk ma cztery wyprowadzenia, po dwa połączone ze sobą wewnętrznie. Naciśnięcie przycisku powoduje zwarcie obu par styków. Z tego powodu elektrycznie wystarczy wykorzystać dowolne dwa przeciwległe piny.
- Orientacja montażu mikroprzycisku na płytce PCB nie ma znaczenia elektrycznego, o ile jego wyprowadzenia pasują do otworów w płytce. Ważne jest jedynie prawidłowe dopasowanie mechaniczne do footprintu.
- Włóż mikroprzycisk do płytki tak, aby wszystkie cztery wyprowadzenia swobodnie przeszły przez otwory. Element powinien przylegać płasko do powierzchni PCB.
- Wyprowadzenia mikroprzycisku są specjalnie wyprofilowane w taki sposób, aby po włożeniu do płytki były w niej sprężyste trzymane. Dzięki temu element sam utrzymuje się w otworach PCB i nie ma potrzeby wyginania wyprowadzeń ani przytrzymywania przycisku podczas lutowania.
- Mimo wszystko, nie zaszkodzi postąpić w sposób konwencjonalny: najpierw przylutuj jedno wyprowadzenie w narożniku. Pozwoli to ustabilizować położenie przycisku i w razie potrzeby skorygować jego ustawienie względem płytki.
- Po upewnieniu się, że mikroprzycisk równo przylega do PCB, przylutuj pozostałe trzy wyprowadzenia.

Zapewni to dobrą stabilność mechaniczną elementu.

- Prawidłowe luty powinny być gładkie, błyszczące i mieć kształt niewielkiego stożka. W przypadku wątpliwości warto zajrzeć do sekcji *Lutowanie komponentów przewlekanych do płytki drukowanej*, gdzie opisano poprawną technikę i zasady bezpieczeństwa.
- Po zakończeniu montażu sprawdź działanie przycisku. Powinien dawać wyraźne „kliknięcie” i sprężyste wracać do pozycji wyjściowej po puszczeniu.
- W przypadku mikroprzycisków nie ma potrzeby przycinania ich wyprowadzeń. Po przylutowaniu pozostaw je w oryginalnej długości.

Montaż przełącznika zasilania

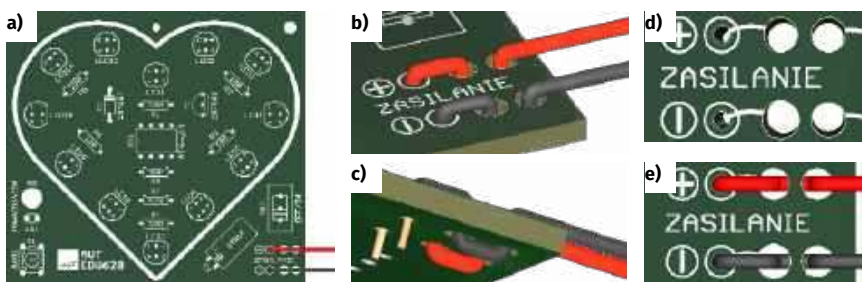
Zgodnie z informacjami z listy elementów zamontuj przełącznik SW1 na wskazanej lokalizacji. Przed przystąpieniem do tej czynności zapoznaj się ze wskazówkami poniżej.

- Przełącznik zasilania SW1 to element, który przełącza połączenie pomiędzy pinem środkowym a jednym ze skrajnych, zależnie od położenia hebelka. W tego typu konstrukcji kierunek montażu nie ma żadnego znaczenia – niezależnie od tego, jak zostanie obrócony, będzie działał prawidłowo.
- Obrys na płytce PCB może zawierać dodatkowe linie symbolizujące położenie hebelka, ale służą one wyłącznie temu, by łatwo rozpoznać miejsce montażu. Nie są to oznaczenia biegunowości ani wymaganej orientacji.

- Włóż przełącznik do płytki tak, aby wszystkie trzy jego wyprowadzenia swobodnie przeszły na drugą stronę PCB. Piny przełącznika są sztywne i nie nadają się do wyginania, dlatego nie należy ich odchyłać w celu stabilizacji elementu.
- Ponieważ wyprowadzenia są sztywne, przełącznik trzeba przytrzymać podczas lutowania – można to zrobić ręką albo poprosić o pomoc kolegę lub opiekuna.
- Najpierw przylutuj środkowy pin przełącznika. Ten pojedynczy punkt lutowniczy pozwala ustabilizować komponent i kontrolować jego położenie względem płytki.
- Po upewnieniu się, że przełącznik dobrze przylega do PCB, przylutuj pozostałe dwa wyprowadzenia. Dzięki temu unikniesz sytuacji, w której element zostanie przylutowany pod kątem lub z przerwą pomiędzy obudową a powierzchnią płytki.
- Luty powinny być czyste, błyszczące i solidne. Jeśli nie masz pewności, czy lutujesz poprawnie, zajrzyj do sekcji *Lutowanie komponentów przewlekanych do płytki drukowanej*, gdzie opisano zarówno technikę, jak i zasady bezpieczeństwa.
- Po zakończeniu montażu sprawdź mechaniczne działanie hebelka. Przełącznik powinien poruszać się lekko i wyraźnie wskazywać w dwie pozycje pracy.
- W przypadku tego typu przełącznika nie ma potrzeby przycinania



Rysunek 12. a) lokalizacja termistora NTC na płytce; b), c), d) poprawny montaż tego elementu



Rysunek 13. a) lokalizacja przewodów na płytce; b), c), d), e) poprawny ich montaż

jego wyprowadzeń. Po przylutowaniu pozostaw je w oryginalnej długości.

Montaż termistora NTC

Zgodnie z informacjami z listy elementów zamontuj termistor NTC na wskazanej lokalizacji. Przed przystąpieniem do tej czynności zapoznaj się ze wskazówkami poniżej.

- Termistor to element, którego rezystancja zmienia się wraz z temperaturą. Wyróżnia się dwa podstawowe rodzaje termistorów: NTC i PTC, różniące się kierunkiem zmian rezystancji pod wpływem temperatury.
- Termistor NTC (Negative Temperature Coefficient) charakteryzuje się rezystancją malejącą wraz ze wzrostem temperatury, natomiast termistor PTC (Positive Temperature Coefficient) – rezystancją rosnącą wraz ze wzrostem temperatury. Oba typy stosowane są między innymi do pomiaru temperatury, jej kontroli, zabezpieczeń termicznych oraz kompensacji zmian parametrów układu.
- Termistory nie mają biegunowości – działają tak samo niezależnie od kierunku włączenia do obwodu. Podczas montażu nie trzeba więc zwracać uwagi na orientację wyprowadzeń, a jedynie na właściwe miejsce na płytce drukowanej.
- Choć termistory NTC i PTC są zazwyczaj niewielkich rozmiarów, często mają cienkie i delikatne wyprowadzenia oraz bywają montowane w sposób nietypowy, na przykład odsunięte od płytki lub w pobliżu źródeł ciepła. Z tego powodu bardzo często lutuje się je na późniejszym etapie montażu, po przylutowaniu większości pozostałych elementów.
- Przed wlutowaniem termistora sprawdź w liście elementów, jaki typ (NTC czy PTC) oraz jaka wartość są przewidziane w danej lokalizacji (np. NTC 10 kΩ). W razie wątpliwości upewnij się, że sięgasz po właściwy element.
- Jeśli to możliwe, zmierz rezystancję termistora za pomocą multimetru ustawionego w tryb omomierza. Pamiętaj, że wynik pomiaru zależy od temperatury otoczenia oraz od tego, czy element nie jest ogrzany, na przykład dotykiem dłoni. W temperaturze odniesienia 25°C pomiar powinien wskazać rezystancję nominalną.
- Wyprowadzenia termistora zegnij ostrożnie i z wycuciem, unikając wielokrotnego doginania w tym samym miejscu, co mogłoby doprowadzić do ich osłabienia lub złamania.
- Umieść termistor w płytce drukowanej w miejscu oznaczonym odpowiednim desygnatorem (patrz rysunek), zwracając uwagę, aby nie naprężyć jego wyprowadzeń.
- Postępuj zgodnie ze wskazaniem instrukcji montażu, jeśli przewiduje ona określony sposób lub miejsce instalacji tego elementu, na przykład wymaganą odległość od płytki.
- Przylutuj termistor możliwie krótkim czasem grzania, aby nie przegrzać elementu. W razie potrzeby zapoznaj się z sekcją *Lutowanie komponentów przewlekanych do płytki drukowanej*.
- Po zakończeniu lutowania odetnij nadmiar wyprowadzeń obcinaczkami. Wskazówki dotyczące bezpiecznego przycinania znajdziesz w sekcji *Przycinanie nadmiaru wyprowadzeń*.

Montaż koszyeczka baterii

Na koniec przylutuj kabelki koszyeczka na trzy baterie AA. Przed przystąpieniem do tej czynności zapoznaj się ze wskazówkami poniżej (rysunek 13).

Pierwsze podłączenie zasilania do płytki

Zanim włożymy drogi i kłopotliwy w nabyciu, zaprogramowany mikrokontroler do podstawki, warto na chwilę podłączyć zasilanie do jeszcze „nieukończonych” płytki i sprawdzić, czy napięcie na złączu zasilania lub na bateriach nie spada do zera



Rysunek 14. Pomiar napięcia na kondensatorze C1 po ustawieniu SW1 w pozycji ON



Rysunek 15. Pomiar napięcia zasilającego mikrokontroler na odpowiednich wyprowadzeniach podstawki po ustawieniu przełącznika SW1 w pozycję ON. Użyj schematu ideowego by ustalić właściwe wyprowadzenia lub skorzystaj z rysunku powyżej

– taki spadek oznaczałby błąd montażowy lub zwarcie w obwodzie zasilania. **Przed przystąpieniem do tej czynności przeczytaj koniecznie sekcję Zagadnienia BHP związane z uruchamianiem zmontowanego układu.** Na początku ustawiamy włącznik w pozycji OFF i mierzymy napięcie na wejściu, czyli na przewodach baterii lub zasilacza. Jeśli pojawia się tam spodziewana wartość (napięcie znamionowe zasilacza lub kompletu baterii), oznacza to, że źródło zasilania działa poprawnie. Następnie przełączamy włącznik zasilania oznaczony na płytce desygnatorem SW1 w pozycję ON i sprawdzamy, czy napięcie nie znika. Na tym etapie napięcie za przełącznikiem można zmierzyć bezpośrednio na wyprowadzeniach kondensatora C1 (rysunek 14).

Jego zanik wskazywałby na usterkę montażową zbudowanego układu, na przykład przypadkowe zwarcie powstałe przez połączenie cyną dwóch sąsiednich pól lutowniczych. W namierzeniu takich pomyłek bardzo pomaga schemat montażowy, na którym dokładnie widać przebieg ścieżek, rozmieszczenie padów oraz to, które połączenia powinny istnieć, a których być nie powinno. Po odnalezieniu usterki należy ją oczywiście wyeliminować. Gdy napięcie wejściowe jest prawidłowe i stabilne, upewniamy się, że po ustawieniu włącznika w pozycję ON pojawia się ono również na odpowiednich pinach

Zagadnienia BHP związane z uruchamianiem zmontowanego układu

- Przed podłączeniem zasilania bezwzględnie załóż okulary ochronne w celu ochrony oczu przed możliwym wybuchem nieprawidłowo zamontowanych elementów spolaryzowanych, takich jak kondensatory elektrolityczne, układy scalone, tranzystory czy diody.
- Upewnij się, że używasz zasilacza lub baterii o właściwym napięciu i prawidłowej polaryzacji – błędne podłączenie zasilania może spowodować gwałtowne uszkodzenie lub wybuch elementów w uruchamianym układzie.
- Upewnij się, że w układzie nie ma zwarć ani luźnych metalowych ścinków na płytce. Zwarcie może spowodować nagłe przegrzanie i wybuch elementów.
- Nie pochylaj twarzy nad układem podczas pierwszego włączania zasilania. Zachowaj rozsądną odległość, aby zminimalizować ryzyko urazu w przypadku awarii elementu.
- Jeśli po włączeniu zasilania pojawi się dym, trzask, zapach spalenizny lub nagłe nagrzanie któregoś z elementów – natychmiast odłącz zasilanie, a następnie – przy odłączonym zasilaniu – znajdź i usuń usterkę.

podstawki mikrokontrolera. Numery pinów zasilających odczytujemy ze schematu ideowego.

Takie testowe podłączenie zasilania oraz pomiary wykonane na źródle zasilania i na odpowiednich pinach podstawki pozwalają wstępnie zweryfikować poprawność montażu oraz uchronić najdroższe i najbardziej wrażliwe elementy – układy scalone, w tym mikrokontroler – przed uszkodzeniem.

Montaż układu scalonego w podstawce

Jeśli woltomierz podczas pomiaru napięcia na źródle zasilania oraz na odpowiednich pinach podstawki wskazał właściwe wartości (tu około 4,5 V) możesz odłączyć baterie, wyłączyć przełącznik SW1 i zamontować mikrokontroler w podstawce. W przeciwnym wypadku, musisz odłączyć zasilanie, odnaleźć i naprawić usterkę montażową.

Włożenie układu scalonego do podstawki wydaje się łatwe, ale trzeba przy tym zachować uwagę i ostrożność. Ważne jest, aby układ był skierowany we właściwą stronę (co wytłumaczę za moment) oraz by wszystkie jego nóżki trafiły dokładnie w otwory podstawki. Nóżki nie mogą się wygiąć ani tym bardziej złamać. A gdyby jednak któraś się urwała, nic straconego – można ją zastąpić kawałkiem odciętego pinu z innego elementu już przyłutowanego do płytki.

Drugą, obok ostrożności podczas montażu sprawą, o jaką należy zadbać, to właściwy kierunek montażu układu scalonego w podstawce. W tym celu należy

przypilnować, by kropka lub wycięcie na układzie scalonym, wskazujące kierunek montażu, pokrywało się z pozostałymi znacznikami w podstawce oraz na warstwie opisowej PCB (**rysunek 16**). Gdy lokalizacja znacznika na układzie scalonym zgadza się z pozostałymi, można przystąpić do wciśnięcia układu w podstawkę.

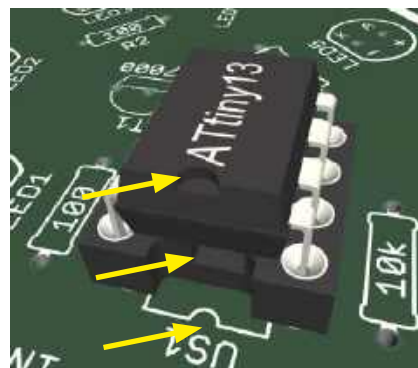
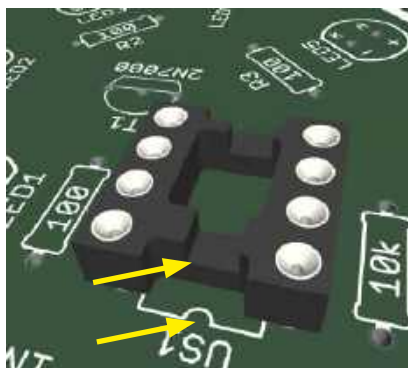
Lutowanie komponentów przewlekanych do płytki drukowanej

Montaż elementów przewlekanych na płytce drukowanej zaczyna się od prawidłowego i stabilnego umieszczenia komponentów w otworach PCB. Najpierw przygotuj element, który chcesz zamontować. W przypadku komponentów spolaryzowanych, takich jak diody, tranzystory i podstawki pod układy scalone należy pamiętać o zachowaniu właściwego kierunku ich osadzenia. W przypadku elementów o dwóch wyprowadzeniach, takich jak rezystory czy diody, należy

delikatnie wygiąć ich nóżki tak, aby można było swobodnie włożyć je w przeznaczone dla nich otwory (patrz załączone rysunki). Po wsunięciu komponentu i docięnięciu jego korpusu do powierzchni płytki warto lekko odchylić wyprowadzenia na zewnątrz. Dzięki temu element nie wypadnie po odwróceniu płytki. Przy komponentach o trzech nóżkach, na przykład tranzystorach, postępuje się podobnie, jednak odgina się wyłącznie skrajne wyprowadzenia, pozostawiając środkowe w naturalnej pozycji. Ułatwia to ewentualny demontaż elementu w przyszłości i zmniejsza ryzyko jego przypadkowego uszkodzenia.

Szczególną uwagę należy zwrócić na montaż podstawek pod układy scalone, ponieważ wygląda on nieco inaczej niż w przypadku zwykłych elementów. Podstawki precyzyjne mają sztywne, okrągłe piny, których nie da się odgiąć, dlatego po włożeniu do płytki podstawka wymaga przytrzymania. Najłatwiej zrobić to, dociskając podstawkę palcem lub opierając płytkę o blat tak, aby podstawka nie mogła wypaść z otworów. Następnie przylutowuje się po jednym skrajnym wyprowadzeniu w dwóch przeciwnych rogach. Gdy podstawka jest wstępnie unieruchomiona, należy sprawdzić, czy całą powierzchnią przylega do płytki. Jeśli któryś róg odstaje, wystarczy ponownie rozgrzać dany lut i delikatnie docisnąć podstawkę; po ostygnięciu powinna ona leżeć idealnie równo.

W zestawach do samodzielnego montażu najczęściej spotyka się jednak podstawki zwykłe, których wyprowadzenia wykonane są z cienkiej, sprężystej blachy. Takie piny można po włożeniu podstawki do płytki lekko odgiąć na zewnątrz, co wystarczająco unieruchomi podstawkę i zapobiegnie jej wypadnięciu podczas odwracania płytki. Wymaga to jednak szczególnej



Rysunek 16. Przed zamontowaniem mikrokontrolera w podstawce należy upewnić się, że znaczniki kierunku montażu – na płytce, w podstawce i na mikrokontrolerze – znajdują się w tej samej pozycji



Rysunek 17. Wskazówki dotyczące sposobu lutowania

ostrożności. Delikatne, płaskie wyprowadzenia bardzo łatwo się wyginają i mogą zamiast do otworu trafić pod podstawkę, chowając się między nią a płytka. To jeden z najbardziej zdradliwych błędów montażowych – trudny do zauważenia i równie trudny do naprawy – dlatego przed rozpoczęciem lutowania koniecznie sprawdź,

czy absolutnie każdy pin wszedł do właściwego otworu. Dopiero gdy wszystkie nóżki są poprawnie osadzone, można rozpocząć lutowanie właściwe.

Rozgrzanym do około 350°C grotem lutownicy łatwo się poparzyć, dlatego nigdy nie należy dotykać metalowej części urządzenia. Do trzymania służy plastikowa,

zwykle gumowana rękojeść, która zapewnia że nie ślizga się ona i nie przesuwają w dłoni. Lutownicę trzymaj w tej ręce, którą piszesz, tak jak trzyma się ołówek – stabilnie, ale bez zbędnego nacisku. Podczas lutowania grot powinien jednocześnie dotykać pola lutowniczego i wyprowadzenia elementu. Tylko wtedy cyna prawidłowo rozpułynie się między nimi. Po krótkim nagraniu, trwającym zwykle dwie do trzech sekund, przykładamy cynę do miejsca lutowania, a nie do grotu. Roztopiona cyna powinna otoczyć nóżkę elementu i utworzyć mały, lśniący stożek. Kiedy ten kształt się pojawi, należy odsunąć drut cyny, a po chwili także lutownicę. Poprawnie wykonany lut jest błyszczący, równy i trwały, a wszystkie elementy pozostają stabilnie zamocowane. Jakość wykonanego połączenia możesz zweryfikować wspomagając się **rysunkiem 17**. Dzięki spokojnej pracy i starannemu przygotowaniu nawet pierwsze lutowanie może zakończyć się naprawdę profesjonalnym efektem.

Zagadnienia BHP związane z lutowaniem

1. Gorący grot to ryzyko poparzenia (dla Ciebie i innych)

- Rozgrzany grot lutownicy ma około 350°C – nie dotykaj go i nie i nie zbliżaj do skóry.
- Odkładaj lutownicę wyłącznie na stabilną podstawkę, aby nikt przypadkiem jej nie dotknął.
- Nie machaj lutownicą i nie przechodź z nią nad innymi osobami.

2. Ochrona oczu (wskazana zwłaszcza w pracy grupowej)

- Topnik zawarty w cynie potrafi pryskać, co bywa odczuwalne na dłoniach. Okulary ochronne zmniejszają ryzyko kontaktu takich odprysków z okolicą oczu.
- Okulary ochronne chronią przed przypadkowym zbliżeniem gorących elementów lutownicy do twarzy.
- W pracy grupowej okulary dodatkowo zmniejszają ryzyko urazu spowodowanego czyjąś nieuwagą.
- W pracy indywidualnej, jeśli okulary bardzo ograniczają widoczność, można z nich zrezygnować – pod warunkiem, że pracujesz ostrożnie i nie zbliżasz twarzy do lutowanych elementów, dzięki czemu ograniczasz ryzyko kontaktu z pryskającym topnikiem.

3. Opary topnika

- Topnik w rdzeniu bezołowiowej cyny nadal wydzielą opary – nie wdychaj ich bezpośrednio.
- Stosuj pochłaniacz oparów lub zapewnij przewiew, aby chronić drogi oddechowe swoje i osób w pobliżu.

4. Ochrona przed pożarem

- Odsuń papier, chusteczki, folie i inne łatwopalne przedmioty od lutownicy.
- Nie zostawiaj lutownicy bez kontroli – ktoś może ją strącić lub dotknąć.
- Upewnij się, że kabel lutownicy nie jest uszkodzony oraz nie ociera się o gorące elementy.

5. Procedury po zakończeniu pracy

- Odłącz lutownicę od prądu i pozostaw do ostygnięcia w miejscu niedostępnym dla innych.
- Umyj ręce – usuwa się pozostałości topnika i metali (mimo braku ołowiu to wciąż dobra praktyka higieniczna).

6. W razie wypadku

- Poparzenie chłodzi zimną bieżącą wodą przez co najmniej 10 minut.
- Podejrzanie pożaru → natychmiast odłącz lutownicę i oddal osoby od stanowiska.
- W każdym wypadku natychmiast powiadom opiekuna lub inną dorosłą osobę.

Montaż przyjazny naprawom

Komponenty przewlekane (THT) warto lutować do płytki w taki sposób, aby w razie potrzeby można je było łatwo wymienić. Zawsze dobrze mieć z tyłu głowy, że któryś z elementów może się kiedyś uszkodzić albo – już na etapie montażu – okazać się, że użyto niewłaściwego komponentu i konieczna będzie jego wymiana. Od tego, jak zostanie on zamontowany, zależy, czy późniejsza naprawa okaże się udręką, czy raczej prostą, szybką czynnością.

Zaginanie wyprowadzeń elementów przeznaczonych do montażu przewlekane wydaje się kuszące – stabilizuje je to w otworach i zapobiega wypadaniu przed lutowaniem. Niestety, taka metoda znacznie utrudnia demontaż uszkodzonych lub błędnie zamontowanych komponentów. Szczególnie kłopotliwe bywa wylutowywanie diod LED z twardymi nóżkami zagiętymi pod kątem: bardzo łatwo wtedy uszkodzić pole lutownicze na płycie, a czasem też samą diodę.

Aby tego uniknąć, po umieszczeniu elementu w otworach warto przed przylutowaniem jednego z wcześniej zagiętych pinów wyprostować go – ustawić prostopadłe do płytki – i dopiero wtedy przylutować. Tę czynność należy powtórzyć dla każdego wyprowadzenia. Nie trzeba obawiać się, że element wypadnie z płytki: po jej odwróceniu komponent opiera się o blat i jest dodatkowo trzymany przez pozostałe,

Zagadnienia BHP związane z przycinaniem nadmiaru wyprowadzeń

1. Postępuj ostrożnie z ostrymi narzędziami

- Używaj tylko sprawnych cząstków do elektroniki.
- Palce trzymaj z dala od ostrza.
- Cząstków nie podawaj ostrzem w stronę drugiej osoby.

2. Zapobieganie odpryskom (najważniejsza zasada)

- Odcinane wyprowadzenie zawsze obowiązkowo chwytaj drugą ręką – wtedy nie odskoczy i nikogo nie zrani.
- Jeśli drucik jest zbyt krótki, by go chwycić, ustaw drugą ręką tak, aby zastąpiła tor ewentualnego odprysku (z zachowaniem bezpiecznej odległości od ostrza).
- Nie pochylaj twarzy nad miejscem cięcia.

3. Ochrona oczu (obowiązkowa)

- Okulary ochronne są absolutnie obowiązkowe – zwłaszcza podczas pracy grupowej, gdzie odprysk może pochodzić nie tylko z Twojego stanowiska.

4. Porządek i bezpieczeństwo na stanowisku

- Pracuj nad stołem i odkładaj ścinki w jedno miejsce.
- Po pracy sprawdź blat i wyrzuć wszystkie metalowe kawałki do kosza na śmieci.

5. W razie wypadku

- Skaleczenie przemyj wodą, zdezynfekuj i załóż plaster.
- Jeśli ściniek trafi do oka – nie pocieraj, nie uciskaj powieki i nie próbuj go wyjąć. Utrzymaj oko spokojnie i natychmiast powiadom opiekuna.
- O każdym urazie natychmiast poinformuj dorosłego.

Zabezpieczenie przed odwrotną polaryzacją

Dioda **D1** (1N4007) została zastosowana jako zabezpieczenie przed odwrotną polaryzacją zasilania, jednak nie w klasycznym układzie szeregowym, lecz w konfiguracji równoległej, pomiędzy dodatnią szyną zasilania a masą. Przy poprawnym podłączeniu baterii dioda jest spolaryzowana zaporowo i pozostaje całkowicie nieaktywna, nie wpływając na napięcie zasilania układu ani na jego sprawność.

Zastosowano takie rozwiązanie świadomie, aby uniknąć spadku napięcia na diodzie, który w przypadku włączenia szeregowego wynosiłby około 0,7 V. W prostych, bateryjnych układach zasilanych z jednego ogniwa nawet tak niewielki spadek napięcia może mieć istotne znaczenie – skracać czas pracy, zmniejszać jasność diod LED lub prowadzić do nieprawidłowej pracy mikrokontrolera przy częściowo rozładowanej baterii.

W przypadku odwrotnego podłączenia zasilania dioda D1 przewodzi i praktycznie zwiera baterię, powodując gwałtowny spadek napięcia na zaciskach układu. Dzięki temu napięcie docierające do mikrokontrolera i pozostałych elementów nie osiąga wartości mogących spowodować ich uszkodzenie. Takie zabezpieczenie jest skuteczne pod warunkiem użycia źródła zasilania o ograniczonej wydajności prądowej, jakim jest bateria lub akumulator, co zostało spełnione w tym projekcie. Zastosowanie diody w konfiguracji równoległej jest więc kompromisem pomiędzy skuteczną ochroną zasilania w normalnych warunkach pracy. Rozwiązanie to jest szczególnie często spotykane w prostych projektach edukacyjnych i urządzeniach zasilanych bateryjnie, gdzie liczy się maksymalne wykorzystanie dostępnego napięcia.

Filtracja zasilania

Kondensator elektrolityczny **C1** o pojemności 100 μF odpowiada za filtrację napięcia zasilania. Jego zadaniem jest tłumienie wolniejszych wahań napięcia oraz krótkotrwałych spadków wynikających z impulsowego poboru prądu przez mikrokontroler i diody LED. Dzięki temu układ pracuje stabilnie, bez niepożądanych resetów czy zakłóceń generowanych przez obciążenie.

Obwód resetu

Rezystor **R7** (4,7 k Ω) pełni rolę rezystora podciągającego linię RESET

nadal zagięte lub już przylutowane nóżki. W przypadku elementów wielonóżkowych stabilizacja jest jeszcze lepsza, ponieważ jednocześnie utrzymuje je w miejscu kilka wyprowadzeń.

Taka technika montażu bardzo ułatwia późniejszą naprawę. Skrócone, niezablokowane zagięciem piny bez trudu przechodzą przez otwory, gdy tylko ponownie rozgrzejemy lut. Dzięki temu wymiana komponentu staje się szybka, bezpieczna i nie grozi uszkodzeniem płytki.

Przycinanie nadmiaru wyprowadzeń

Nadmiar wyprowadzeń uprzednio przylutowanych komponentów do montażu przewlekane obetnij przy użyciu obcinaczek. Pamiętaj, aby nie ścinać lutu – cięcie wykonuj dopiero za miejscem, w którym kończy się spoina i wystaje samo wyprowadzenie (**rysunek 18**). Pamiętaj, aby nie szarpać ani nie ciągnąć za wyprowadzenie przylutowanego komponentu.



Rysunek 18. Właściwe miejsce cięcia nadmiaru wyprowadzeń po przylutowaniu komponentu. Po wykonaniu cięcia spoina powinna pozostać nienaruszona

Mogłoby to doprowadzić do dość kłopotliwego w naprawie wyrwania pola lutowniczego z płytki i zerwania połączenia komponentu ze ścieżką.

Podsumowanie montażu

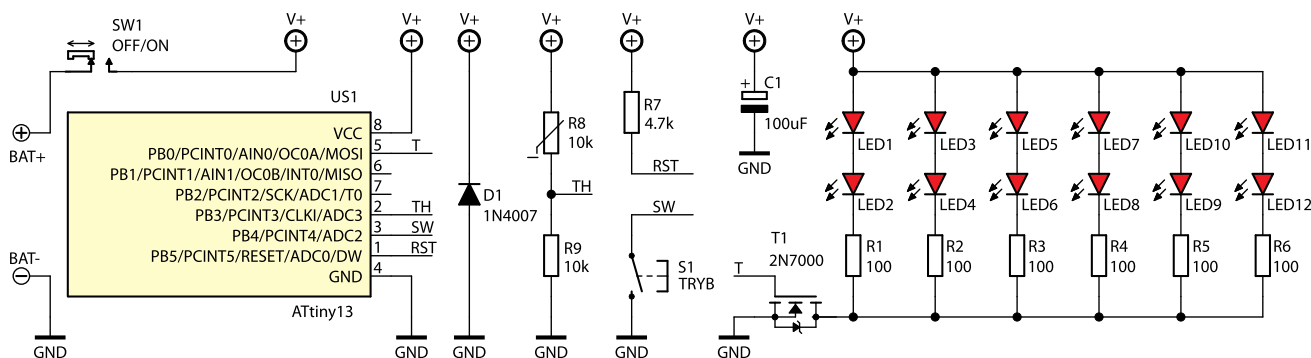
Po ukończeniu montażu sprawdź, proszę, czy wszystkie połączenia lutowane są błyszczące i nie ma zimnych lutów oraz czy żadne sąsiednie pola lutownicze nie są ze sobą błędnie połączone.

Ale właściwie... dlaczego to działa?

Na schemacie ideowym przedstawionym na **rysunku 19** centralnym elementem układu jest uprzednio zaprogramowany przez producenta zestawu mikrokontroler ATtiny13 (US1), który pełni rolę generatora sygnału sterującego pracą diod LED. Układ został zaprojektowany w taki sposób, aby na podstawie temperatury mierzonej przez termistor R8 możliwe było uzyskanie określonych efektów świetlnych realizowanych programowo, bez konieczności stosowania klasycznych generatorów RC czy bramek logicznych.

Włącznik zasilania

Mikrokontroler zasilany jest bezpośrednio z baterii poprzez włącznik **SW1**, który umożliwia całkowite odłączenie zasilania, zapobiegając rozładowaniu baterii w czasie, gdy układ nie jest używany. Napięcie zasilania oznaczone jako V+ rozprowadzane jest do poszczególnych bloków układu.



Rysunek 19. Schemat ideowy układu

mikrokontrolera do dodatniego biegunu zasilania. Zapewnia on poprawny stan logiczny w czasie normalnej pracy układu oraz zapobiega przypadkowemu resetom spowodowanym zakłóceniami. Więcej informacji dotyczących układów resetu można znaleźć we wcześniejszych odcinkach naszego cyklu. Kilka zdań na ten temat pojawiło się w każdym spotkaniu, w którym opisywaliśmy układ wykorzystujący mikrokontroler:

- Spotkanie nr 4, EdW 4/2024 – UFOLEdek – AVTEDU632
- Spotkanie nr 8, EdW 2/2025 – Wspomagacz wyboru – AVTEDU639
- Spotkanie nr 12, EdW 6/2025 – Przypominacz światło dźwiękowy – AVTEDU646
- Spotkanie nr 14, EdW 8/2025 – TermoEmotek – AVTEDU630

Wejście sterujące

Zastosowanie mikrokontrolera pozwoliło, przy minimalnej liczbie elementów, uzyskać trzy różne efekty świetlne, których tempo można zwiększać poprzez przytrzymanie czujnika temperatury, czyli termistora. Wyboru efektu dokonuje się za pomocą przycisku S1, dołączonego do wyprowadzenia 3 mikrokontrolera (PB4). Sygnał ten na schemacie oznaczono etykietą SW.

Dzielnik napięcia z termistorem

Pomiar temperatury realizowany jest za pomocą termistora NTC R8 (10 kΩ), który wraz z rezystorem R9 (10 kΩ) tworzy dzielnik napięcia dołączony pomiędzy dodatnią szynę zasilania V+ a masę. Napięcie z punktu połączenia obu elementów doprowadzone jest do wejścia analogowego ADC3 mikrokontrolera ATtiny13 (wyprowadzenie 2), opisanego na schemacie etykietą TH. Termistor NTC jest elementem, którego rezystancja maleje wraz ze wzrostem

temperatury. Zmiana jego rezystancji powoduje zmianę napięcia w punkcie TH. Przy niższej temperaturze rezystancja termistora jest większa, a napięcie mierzone przez mikrokontroler przyjmuje wartość bliższą dodatniemu biegunowi zasilania. Wraz ze wzrostem temperatury rezystancja NTC maleje, co powoduje obniżanie się napięcia w punkcie TH. Dobór wartości rezystora R9 równy rezystancji nominalnej termistora przy temperaturze odniesienia (zwykle 25°C) sprawia, że napięcie w punkcie TH znajduje się w środkowej części zakresu napięć zasilania. Zapewnia to dobrą rozdzielczość pomiaru przetwornika analogowo-cyfrowego i umożliwia precyzyjne wykrywanie zmian temperatury w obu kierunkach. Tak zrealizowany dzielnik napięcia jest prostym i skutecznym sposobem pomiaru temperatury w układach zbudowanych z użyciem mikrokontrolerów. Nie wymaga dodatkowych elementów ani skomplikowanej kalibracji, a jego dokładność jest w pełni wystarczająca do sterowania pracą układu i efektami świetlnymi, które mikrokontroler dobiera w zależności od odczytanej temperatury.

Podczas naszych poprzednich spotkań kilkakrotnie omawialiśmy zasadę działania dzielnika napięcia od strony teoretycznej, wraz ze wzorami i przykładowymi obliczeniami. Tym razem postanowiłem ograniczyć wywód do minimum, a zainteresowanych Juniorów odesłać tym samym do naszych juniorskich materiałów archiwalnych (między innymi: *Spotkanie 9*; EdW 3/2025, *Spotkanie 12*; EdW 6/2025).

Sekcja diod LED złączana tranzystorem N-MOSFET

Diody LED1...LED12 zostały zgrupowane w sześć równoległych gałęzi, z których każda zawiera dwie diody połączone szeregowo oraz rezystor ograniczający prąd (R1...R6 o wartości 100 Ω). Rezystory te zabezpieczają diody przed przeciążeniem prądowym

i zapewniają bezpieczne warunki pracy, zgodne z ich parametrami katalogowymi.

Wszystkie gałęzie diod LED są zasilane ze wspólnej szyny V+, natomiast ich załączanie odbywa się poprzez tranzystor T1 typu N-MOSFET (2N7000), który po wystrojeniu dołącza gałęzie z diodami LED do masy układu.

Zastosowanie tranzystora MOSFET zamiast tranzystora bipolarnego pozwala ograniczyć obciążenie wyjścia mikrokontrolera, ponieważ w stanie statycznym bramka MOSFET-a praktycznie nie pobiera prądu sterującego. Choć bramka tranzystora stanowi obciążenie pojemnościowe i przy każdej zmianie stanu logicznego musi zostać naładowana lub rozładowana, w tym układzie – ze względu na stosunkowo niską częstotliwość przełączania oraz niewielką pojemność bramki tranzystora 2N7000 – efekt ten jest niewielki. W praktyce średnie obciążenie wyjścia mikrokontrolera jest znacznie mniejsze niż w przypadku tranzystora bipolarnego, który wymaga stałego doprowadzania prądu bazy. Dodatkowo tranzystor MOSFET powoduje mniejsze straty napięcia w stanie przewodzenia, co ma istotne znaczenie w prostych układach zasilanych z baterii, upraszczając konstrukcję i poprawiając sprawność energetyczną całego układu.

Obwody z diodami LED

Ponieważ sposób zasilania diod LED, a w szczególności dobór odpowiednich wartości rezystorów szeregowych w gałęziach z diodami LED, pojawia się na naszych juniorskich spotkaniach jeszcze częściej niż temat dzielnika rezystancyjnego, pozwolę sobie tym razem również ten temat pominąć.

Zamiast tego ponownie polecam lekturę treści 12. spotkania Juniorów EdW (EdW 6/2025), w szczególności sekcji *Sterowanie diodami LED*. Warto sięgnąć także do śródtytułu *Sekcje diod LED* w treści 17. spotkania

Juniorów (EdW 11/2025), gdzie uwzględniono nie tylko spadki napięć na samych diodach LED i ich rezystorach szeregowych, lecz także na złączach kolektor-emiter sterujących nimi tranzystorów, na szeregowej diodzie prostowniczej zabezpieczającej układ przed odwrotną polaryzacją, a nawet na rezystancji wewnętrznej baterii – czyli na wszystkich elementach rzeczywistego toru przepływu prądu.

Zabawy obliczeniami rozwijają, kształcą, a czasem potrafią wciągnąć niczym gra komputerowa. I wbrew pozorom nie zawsze same liczby są najważniejsze – często cenniejsze okazuje się to, czego doświadczamy w pogoni za nimi.

Nie sposób jednak nie wspomnieć o fakcie, że w porównaniu z tranzystorem bipolarnym tranzystor N-MOSFET w stanie przewodzenia powoduje znacznie mniejszy spadek napięcia. Zamiast stałego spadku rzędu 0,2...0,3 V występuje tu jedynie niewielki spadek wynikający z rezystancji kanału, co poprawia sprawność układu i pozwala lepiej wykorzystać napięcie zasilania. Do tego zagadnienia z pewnością jeszcze niejednokrotnie wrócimy.

Podsumowanie

Tym razem w Twoje ręce trafił zestaw AVTEDU620 – *Bijące serce LED* – niewielki układ, który miga, pulsuje i wygląda jak

prawdziwe serce. Może być sympatycznym prezentem, ozdobą pokoju albo po prostu kolejnym krokiem w nauce lutowania i elektroniki. Ale skoro już mówimy o sercu, warto na chwilę się zatrzymać.

Serce w elektronice bije dokładnie tak, jak przewidywał to programista, odpowiednio pisząc program na mikrokontroler. Bez emocji, bez rozczarowań i bez niespodzianek. W prawdziwym życiu bywa inaczej. Nie każda sympatia kończy się sukcesem, nie każda „iskra” oznacza coś trwałego, a czasem coś, co na początku świeci bardzo jasno, po chwili po prostu gaśnie. Bywa też, że to właśnie „sukces” w perspektywie czasu okazuje się największą pułapką – bo, jak mawiają, miłość bywa ślepa.

Warto pamiętać, że wybory podejmowane pochopnie – także te dotyczące relacji – mogą mieć poważne konsekwencje. Nieodpowiedzialnie dobrany partner potrafi z czasem bardzo skomplikować życie, a nawet je zrujnować. Odpowiedzialność dotyczy jednak nie tylko nas samych. W relacjach zawsze pojawiają się też inni – dzieci, rodzina, bliscy – którzy ponoszą skutki decyzji dorosłych, nawet jeśli nie mieli na nie żadnego wpływu.

Nie każdemu życie w parze – jakiegokolwiek – jest pisane i nie ma w tym nic ani złego, ani nadzwyczajnego. Czasem rozsądniej jest zwolnić, dać sobie czas

i zadbać o własne bezpieczeństwo emocjonalne, zamiast działać pod wpływem impulsu czy presji. Elektronika uczy cierpliwości, dokładności i myślenia krok po kroku – i dokładnie te same cechy przydają się w życiu.

Warto też pamiętać, że żyjemy w czasach, w których technologia oferuje potężne narzędzia. Wiedza, która kiedyś była rozproszona w dziesiątkach książek i specjalistycznych publikacji, dziś jest na wyciągnięcie ręki – dostępna dla każdego i często zupełnie za darmo. Jeszcze niedawno można było tkwić latami w nierozpoznanych mechanizmach, nie rozumiejąc, co właściwie dzieje się w relacji. Dziś rzetelne informacje potrafią pomóc nazwać brak szacunku, presję czy manipulację, a tym samym szybciej zorientować się, że coś przestaje działać tak, jak powinno – albo że nigdy tak nie działało.

W elektronice znajomość podstaw pozwala uniknąć błędów, zanim doprowadzą do uszkodzenia układu – w życiu działa to bardzo podobnie.

Mamy nadzieję, że ten zestaw sprawił Ci radość, ale też stał się pretekstem do krótkiej refleksji. A za miesiąc wrócimy z kolejnym projektem, który – jak zawsze – połączy naukę z dobrą zabawą. Do zobaczenia w następnym numerze! ■

Mariusz Ciszewski

REKLAMA

przejrzyj i kupisz na stronie
www.ulubionykiosk.pl

świat radio
Magazyn wszystkich użytkowników eteru
KROTKOFALARSTWO CB RADIOTECHNIKA

Kalendarz krajowych zawodów krotkofalowych
świat radio
Magazyn wszystkich użytkowników eteru
KROTKOFALARSTWO CB RADIOTECHNIKA

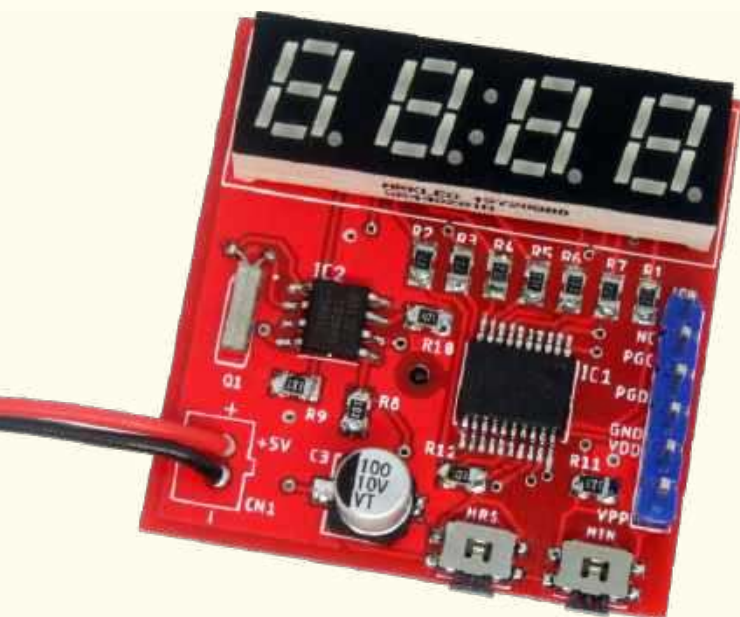
Transceivery retro
Drake TR-4

www.elportal.pl

eprasa.pl fe6a48ff33

Elektronika dla Wszystkich

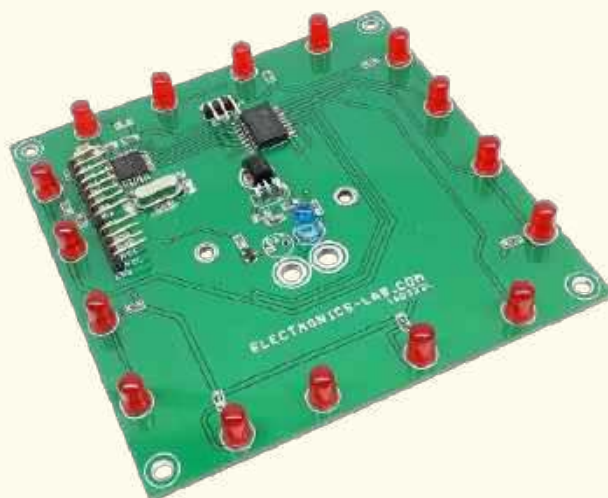
89



7-segmentowy mini zegar wykorzystujący PIC16F628A i DS1307 RTC

Jest to minimalistyczny i niewielki zegar oparty na mikrokontrolerze PIC16F628A i układzie scalonym DS1307 RTC. Wyświetla on czas wyłącznie na niewielkim 7-segmentowym wyświetlaczu składającym się z 4 znaków. Wykorzystano wyświetlacz 0,28" SR440281N RED ze wspólną katodą, ale można również użyć innych wyświetlaczy, takich jak 0,56" Kingbright CC56-21SRWA.

<https://youtu.be/ICekYsTSowC>



Światło LED oparte na czujniku zbliżeniowym

Prezentowany projekt to lampka LED oparta na czujniku zbliżeniowym na podczerwień, zaprojektowana głównie do interaktywnych stolików kawowych, ale może być również używana jako pojedyncza interaktywna lampka. Stolik można wykonać z wielu takich płytek drukowanych. Można użyć wielu płytek, jak pokazano na poniższym obrazku, każda płytka zapali się, gdy nad czujnikiem zostanie umieszczony przedmiot. Zapalą się tylko diody LED znajdujące się w obszarze czujnika. Projekt jest kompatybilny z Arduino.

<https://youtu.be/UP53H-1hqwY>
<https://youtu.be/hokVeo4vOcU>

Niektóre projekty aktualnie dostępne tylko dla prenumeratorów EdW w rubryce **DIY PLUS** na www.elportal.pl:

1. OLEDUINO – wyświetlacz OLED kompatybilny z Arduino
2. Inteligentny regulator lutownicy – precyzyjny regulator grzałki
3. Wskaźnik poziomu paliwa z wyświetlaczem OLED
4. Bezprzewodowy odbiornik wilgotności i temperatury
5. Przerwania zewnętrzne (sprzętowe) i przerwania zegara w MicroPython
6. Laserowy czujnik odległości z wyświetlaczem OLED i RP2040
7. Inklinometr z 17-segmentowym wyświetlaczem słupkowym
8. Izolowany repeater USB – USB 2.0
9. Knight Rider Light – 16 diod LED dużej mocy (kompatybilny z Arduino)
10. Dźwięk do kolorowych efektów świetlnych (kompatybilny z Arduino)
11. Nowy i ulepszony licznik Geigera – teraz z Wi-Fi!
12. Detektor zalania
13. Lampka nastrojowa LED o dużej mocy
14. Kontroler dzwonów kościelnych
15. Arduino Nano – włączanie/wyłączanie urządzeń za pomocą pilota na podczerwień (dwa kanały)
16. Lampka sufitowa LED z czujnikiem ruchu PIR – kompatybilna z Arduino
17. Inteligentny ściemniacz LED z Bluetooth – 4-kanałowy włącznik/wyłącznik Bluetooth
18. Czterokanałowy izolator cyfrowy, wzmacniony, szybki, o niskim poborze mocy
19. Sterowanie prędkością, kierunkiem i zatrzymaniem silnika DC z modułem RF NRF24L01
20. Nadajnik zdalnego sterowania z pojedynczym joystickiem wykorzystujący NRF24L01
21. 8-kanałowy zdalny nadajnik RF z protokołami: Holtek i szeregowym
22. 8-kanałowy zdalny odbiornik RF z protokołami: Holtek i szeregowym
23. Pojemnościowy czujnik wilgotności do konwertera wyjścia analogowego
24. Mostek H dla wysokiej mocy szczotkowego silnika prądu stałego z czujnikiem prądu
25. Przetwornica DC-DC buck 12...75 V na 10 V na wyjściu
26. Czujnik prądu low-side 10 µA...10 mA
27. Kontroler ramienia robota z bezprzewodowym pilotem PS3
28. Termiczny czujnik masowego przepływu powietrza – anemometr statotemperaturowy
29. Precyzyjny wzmacniacz transimpedancjowy z przełączanym integratorem
30. Wysokowydajny monofoniczny wzmacniacz audio klasy D o mocy 20 W
31. Kontroler pełnego mostka z przesunięciem fazowym i prostowaniem synchronicznym wykorzystujący UCC28950
32. Monitorowanie poziomu cieczy za pomocą czujnika ciśnienia – wyświetlacz słupkowy
33. Sterowanie silnikiem DC za pomocą joysticka
34. 16-kanałowy sterownik serwomechanizmów RC z interfejsem I²C
35. Programowalny kondycjoner sygnału z czujnika rezystancyjnego mostkowego
36. Chocinka z Arduino i pikselowymi diodami

Miesięcznik „Elektronika dla Wszystkich” (12 numerów w roku) jest wydawany we współpracy z kilkoma redakcjami zagranicznymi



Wydawnictwo:
 AVTKorporacja Sp. z o.o.
 03-197 Warszawa, ul. Leszczyńska 11
 tel. 22 257 84 99, e-mail: avt@avt.pl

Wydawca:
 Wiesław Marciniak

Redaktor naczelny:
 Mariusz Ciszewski
mariusz.ciszewski@elportal.pl

Adres redakcji:
 03-197 Warszawa, ul. Leszczyńska 11
 e-mail: edw@elportal.pl, www.elportal.pl

Dział reklamy:
 Katarzyna Gugala
katarzyna.gugala@elportal.pl, tel. 22 257 84 64

Szef Pracowni Konstrukcyjnej:
 Jakub Sobański
jakub.sobanski@elportal.pl

Sekretarz redakcji:
 Dariusz Welik
dariusz.welik@elportal.pl

Copyright AVTKorporacja Sp. z o.o., Warszawa, ul. Leszczyńska 11. Projekty publikowane w „Elektronice dla Wszystkich” mogą być wykorzystywane wyłącznie do własnych potrzeb. Korzystanie z tych projektów do innych celów, zwłaszcza do działalności zarobkowej, wymaga zgody redakcji „Elektroniki dla Wszystkich”. Przedruk oraz umieszczanie na stronach internetowych całości lub fragmentów publikacji zamieszczanych w „Elektronice dla Wszystkich” jest dozwolone wyłącznie po uzyskaniu pisemnej zgody redakcji. Redakcja nie odpowiada za treść reklam i ogłoszeń zamieszczanych w „Elektronice dla Wszystkich”.

DTP, redakcja strony internetowej www.elportal.pl:
 MAD Sp. z o.o.

Prenumerata:
 W Wydawnictwie AVT, e-mail: prenumerata@avt.pl
 tel. 22 257 84 22, (godz. 10:00–14:00)
www.ulubionykiosk.pl



FN-SWM10

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



FN-DPOS-350P

Dwukanałowy oscyloskop 350 MHz, FNIRSI DPOS350P



FN-2C53T

Dwukanałowy oscyloskop z multimetrem i generatorem 50 MHz FNIRSI 2C53T

BESTSELLERY sklepu AVT – sklep.avt.pl

Mierniki Testery FNIRSI

Rabat dla Czytelników EdW
przy zakupie podaj kod **EdW2505FN**

Kod ważny do 30.09.2025

-3%

Rabat dla Prenumeratorów EdW
przy zakupie podaj numer prenumeraty

-6%



FN-LCR-ST1

Miernik pęsetowy, tester elementów FNIRSI LCR-ST1



FN-LCR-P1

Tester elementów FNIRSI LCR-P1



FN-HRM10

Tester rezystancji wewnętrznej akumulatorów FNIRSI HRM-10



FN-G1200

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



FN-DWS200-F245

Stacja lutownicza 200 W z kolbą F245, FNIRSI DWS200



FN-1014D

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



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



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



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

BESTSELLERY sklepu AVT – sklep.avt.pl

Trzecia ręka

Rabat dla Czytelników EdW
przy zakupie podaj kod **EdW2505TR**

Kod ważny do 30.09.2025

-3%

Rabat dla Prenumeratorów EdW
przy zakupie podaj numer prenumeraty

-6%



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



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



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



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



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



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