

ELEKTRONIKA PRAKTYCZNA

EP.com.pl

● Międzynarodowy magazyn elektroników konstruktorów ● styczeń ● 1/2025 ●

Tylko Prenumeratorzy

- mają dostęp do artykułów przed ich publikacją w EP na www.ep.com.pl – **EP W TOKU**
- mają dostęp do materiałów dodatkowych, takich jak pliki źródłowe projektów na naszym serwerze **FTP** www.ulubionykiosk.pl/media

inspirujące, użyteczne projekty

- Mikser dyskotekowy trzech źródeł sygnału • Retro wzmacniacz mocy z tranzystorami germanowymi • TRX Ewa 40 m – transceiver QRP początkującego krótkofalowca
- Przedwzmacniacz gramofonowy MM z pasywną korekcją
- Theremin – elektroniczny instrument

podzespoły, sprzęt, aplikacje

- Wyświetlacze dotykowe od Artronic • Małe, ale wariaty: nowoczesne, scalone wzmacniacze audio małej mocy
- Wyświetlacze dotykowe i sensory gestów

tutoriale

- Druk 3D w służbie elektroniki • Internet Rzeczy w pomiarach środowiskowych. Dołączanie rezystancyjnych i pojemnościowych czujników deszczu do modułu Enviro Weather • Źródło prądowe (prawie) idealne • Ograniczenia diody idealnej • Nierozumiana pojemność kabli • Konwersja analogowo-cyfrowa w technice audio

kursy

- Pomiarów charakterystyk częstotliwościowych. Filtry aktywne m.cz. typu Sallen-Key • Czytniki kodów kreskowych w praktyce. Frontem do użytkownika: moduł skanera 1D/2D EM20-80 V2 • Implementacja systemu Linux na platformie STM32MP. Pierwsza aplikacja • Kurs FPGA Lattice. Obsługa monitora VGA

TECHNOLOGIE AUDIO

TEMAT NUMERU



Nowy dział: **AUDIO BEZ TAJEMNIC**



NIE TYLKO WYŚWIETLACZE DOTYKOWE

–20%
NA START
181,40 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 Praktyczną*. Decydując się teraz na roczną prenumeratę drukowaną, otrzymasz nie tylko dostęp do najnowszych wydań, ale i **znakomity start dzięki zniżce 20%** 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 126,90 zł.

Dodatkowa oferta dla prenumeratorów wersji drukowanej: jeśli już subskrybujesz wersję papierową, możesz dokupić równoległe 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:

- artykułów przed ich publikacją w *Elektronice Praktycznej* na www.ep.com.pl – EP W TOKU
- materiałów dodatkowych (takich jak pliki źródłowe projektów) na www.UlubionyKiosk.pl/media

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.



Zacznij korzystać z pełnych zasobów już dziś!

Magia

Dziewiątego kwietnia 1860 roku. Dopiero za 17 lat świat obiegnie wieść o skonstytuowaniu przez Edisona fonografu, czyli (rzekomo) pierwszego urządzenia zdolnego do zapisu dźwięku. A jednak w swojej paryskiej pracowni Édouard-Léon Scott de Martinville – do dziś relatywnie mało znany księgarz i zecer – dokonuje pierwszego w historii zapisu dźwięku z użyciem autorskiego urządzenia zwanego fonautografem. Swój wynalazek – opracowany zresztą z myślą o graficznym zapisie fal akustycznych w celach naukowych – opatentował on trzy lata wcześniej, a działające urządzenie zastosował do wykonania krótkiego nagrania. Twórca fonautografu nie dysponował jednak – podobnie jak i inni wynalazcy w owym czasie – jakąkolwiek techniką pozwalającą na odtworzenie dźwięku. Zapis trafił do paryskiego archiwum i przeleżał w nim 148 lat. Scottowi nie udało się także utrzymać palmy pierwszeństwa: sława trafiła do Edisona, choć w swoich pamiętnikach, spisanych tuż przed śmiercią, wynalazca żalił się, że to jemu (i Francji), a nie Edisonowi i Ameryce, należy się splendor za stworzenie wynalazku. Frustracja Europejczyka była zresztą zrozumiała, wszak fonograf Edisona opierał się na bardzo zbliżonych założeniach konstrukcyjnych co fonautograf.



Rok 2008. Berkeley, około 9000 km od Paryża. Grupa naukowców, posługując się odpowiednim skanerem cyfrowym i specjalnym oprogramowaniem, dokonuje próby rekonstrukcji oryginalnego brzmienia nagranych przez Scotta. I tak do uszu badaczy, a niedługo później całego świata, dociera dźwięk, który dziennikarze „Gazety Wyborczej” słusznie określili mianem „nagrania z zaświatów” – niewyraźne i zaszumione „mruczenie” zostało zidentyfikowane jako fragment ludowej pieśni francuskiej pt. *Au Clair de la Lune*, zaśpiewanej przez... samego wynalazcę! Pierwsze w historii świata nagranie dźwiękowe odtworzono po prawie 150 latach od jego dokonania.

Analogowe metody rejestracji audio przetrwały do dziś i wciąż pozostają „w cenie” – pomimo upowszechnienia nośników cyfrowych oraz sposobów programowego przetwarzania i odtwarzania dźwięku, pocziwe gramofony nadal są chętnie używane przez miłośników dobrego brzmienia. Audiofile mają zresztą „wrodzony” pociąg do technologii trącących myszką, o czym świadczy niegasnące zainteresowanie urządzeniami lampowymi. Współczesny świat audio okazuje się na szczęście bardzo liberalny pod względem ucyfrowienia, więc każdy wybiera taką ścieżkę rozwoju swojego zestawu sprzętowego, która okazuje się najbliższa jego sercu i... portfelowi. A rozrzut cen jest naprawdę spory – niektórym wystarczy zestaw głośnikowy za dwa tysiące złotych, inni za tyle kupują... sam tylko kabel sieciowy.

W styczniowym numerze „Elektroniki Praktycznej” rozpoczynamy zapowiadany w poprzednim miesiącu, całkowicie nowy dział stały o wdzięcznym tytule „Audio bez tajemnic”. Kiedy wpadłem na pomysł utworzenia takiej rubryki na łamach naszego czasopisma, spodziewałem się wprawdzie pozytywnego odzewu ze strony naszych Czytelników, ale... nie przewidziałem, że nastąpi on jeszcze przed jego właściwym uruchomieniem! Tymczasem już w grudniu – tuż po tym, jak do sprzedaży trafił ostatni numer EP datowany na rok 2024 – zaczęliśmy otrzymywać pierwsze maile, a nawet telefony od Czytelników zainteresowanych publikowanymi u nas materiałami czy też proponującymi, abyśmy zajęli się w ramach „Audio bez tajemnic” konkretnymi zagadnieniami, które są dla nich ważne. Dziękuję za ten odzew – nie ukrywam, że taka pozytywna reakcja bardzo mnie cieszy, a kolejne wiadomości, które otrzymuję na redakcyjną skrzynkę e-mailową pokazują, że decyzja była słuszna, ba! – że nasi Czytelnicy czekali na ten ruch i... doczekali się!

Aby nadać uruchomieniu nowego działu stosowną oprawę, naszą styczniową „ramówkę” nasyciliśmy treściami związanymi z audio. Tym razem aż cztery z pięciu projektów to rozmaite urządzenia akustyczne – retro wzmacniacz, przedwzmacniacz z korekcją RIAA, elektroniczny theremin i mikser dyskotekowy. Do tego dochodzi pierwszy artykuł z cyklu poświęconego konwersji analogowo-cyfrowej, poradnik dotyczący wpływu pojemności okablowania na parametry pracy sprzętu oraz przeglądowy materiał, poświęcony nowoczesnym rozwiązaniom scalonym stosowanym w sprzęcie konsumenckim i profesjonalnym. Tak spory rozrzut tematyczny to doskonały dowód na to, że tematyce audio chcemy przyglądać się z różnych stron i odmiennych punktów widzenia – a już teraz nasz zespół redakcyjny w pocie czoła opracowuje kolejne, niezwykle interesujące materiały.

A jakie zagadnienia trafią do naszego działu w przyszłości? To zależy także od Was, drodzy Czytelnicy. Serdecznie zapraszam konstruktorów i serwisantów doświadczonych w różnych obszarach branży elektroakustycznej do współtworzenia rubryki audio – jak zawsze Redakcja EP pozostaje otwarta na współpracę z nowymi Autorami, którzy mają odpowiedni zasób wiedzy i chęć jej przekazania szerokiemu gronu Czytelników.

W najnowszym numerze EP przygotowaliśmy także szereg materiałów spoza tematyki akustycznej. W dziale Projektów publikujemy pierwszą z dwóch części opisu prostego transceivera TRX Ewa 40 m, opartego wyłącznie na elementach dyskretnych: urządzenie, opracowane przez legendę polskiego krótkofalarstwa – Andrzeja Janeczka (SP5AHT), redaktora naczelnego magazynu „Świat Radio” – zostało przygotowane z myślą o początkujących radioamatorach, którzy pragną dogłębnie poznać i zrozumieć działanie sprzętu nadawczo-odbiorczego. Kontynuujemy także cykle poświęcone drukowi 3D i aplikacjom IoT, a w dziale kursów publikujemy kolejne materiały dotyczące obsługi czytników OEM, implementacji systemu Linux na platformie STM32MP, programowania FPGA i pomiarów charakterystyk filtrów.

Zapraszam do lektury, życząc jednocześnie wszelkiej pomyślności i wielu inspiracji w nowym roku 2025!

Przemysław Musz

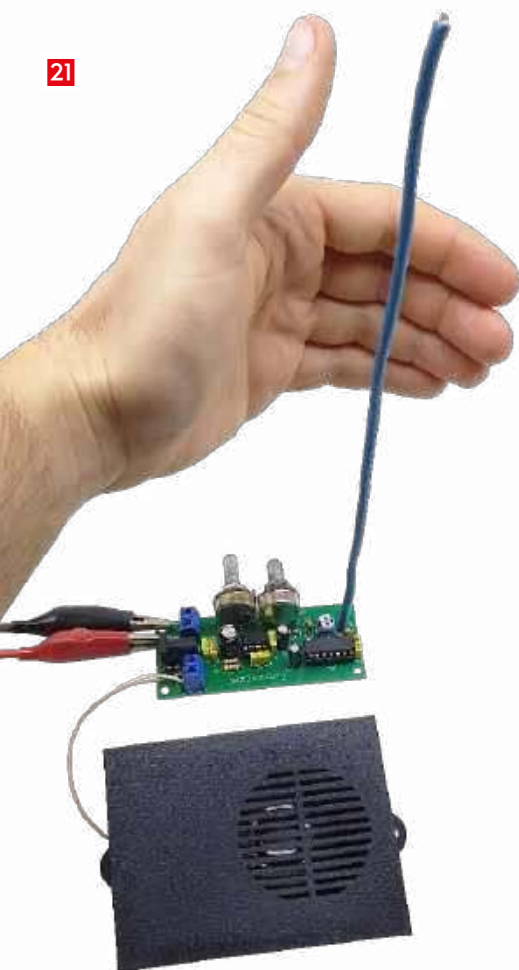
24



18



21



Nie przeocz

Nowe podzespoły	5
Dodaj do obserwowanych	11
Koktajl niusów	92

Projekty

Mikser dyskotekowy trzech źródeł sygnału (1)	12
Retro wzmacniacz mocy z tranzystorami germanowymi	18
TRX Ewa 40 m – transceiver QRP początkującego krótkofalowca (1)	24

Miniprojekty

Przedwzmacniacz gramofonowy MM z pasywną korekcją	16
Theremin – elektroniczny instrument	21

Prezentacje

Wyświetlacz dotykowy od Artronic	62
--	----

Audio bez tajemnic

Nierozumiana pojemność kabli	30
Konwersja analogowo-cyfrowa w technice audio (1)	32

Temat numeru: Technologie audio

Małe, ale wariaty: nowoczesne, scalone wzmacniacze audio małej mocy	38
---	----

Moduły w aplikacjach

Internet Rzeczy w pomiarach środowiskowych (13). Dołączanie rezystancyjnych i pojemnościowych czujników deszczu do modułu Enviro Weather	58
--	----

Technologie wokół elektroniki

Druk 3D w służbie elektroniki (5)	52
---	----

Elektronika w praktyce

Wyświetlacz dotykowy i sensory gestów	63
---	----

Notatnik konstruktora

Źródło prądowe (prawie) idealne	28
Ograniczenia diody idealnej	50

Kursy

Pomiary charakterystyk częstotliwościowych (2).	
Filtry aktywne m.c. typu Sallen-Key	71
Czytniki kodów kreskowych w praktyce (2).	
Frontem do użytkownika: moduł skanera 1D/2D EM20-80 V2	78
Implementacja systemu Linux na platformie STM32MP (2). Pierwsza aplikacja	84
Kurs FPGA Lattice (27). Obsługa monitora VGA	86

Prenumerata	2
Od wydawcy	3
Hity następnego numeru	95

nowe podzespóły

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

Złącza FPC One Action o grubości 0,65 mm

Firma Hirose oferuje niskoprofilowe złącza FPC One Action z serii FH82, których grubość ograniczono do zaledwie 0,65 mm. Są to najmniejsze tego typu złącza z dotychczasowej oferty, mogące znaleźć zastosowania m.in. w motoryzacji, elektronice medycznej, aplikacjach smart home czy urządzeniach przenośnych. Najczęściej służą do łączenia wyświetlaczy bądź modułów czujników z płytą drukowaną.



Określenie „one action” odnosi się do mechanizmu blokady, znacząco ułatwiającego montaż przewodu. W tradycyjnych złączach FPC/FFC instalacja wymaga otwarcia zamka, wsunięcia przewodu FPC/FFC, a następnie zamknięcia blokady. Złącza One Action eliminują te dodatkowe kroki, umożliwiając automatyczne zablokowanie przewodu w momencie jego wsunięcia do złącza. Dzięki temu integracja urządzenia jest szybsza, a także łatwiejsza do zautomatyzowania, co okazuje się szczególnie istotne na nowoczesnych liniach produkcyjnych, na których coraz częściej stosuje się roboty.

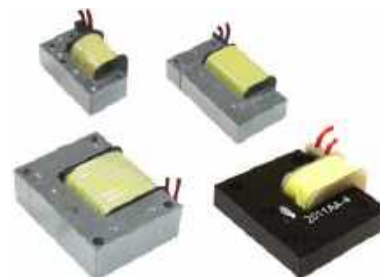
Nowe, niskoprofilowe złącza One Action z serii FH82 są obecnie dostępne w wersji 14-torowej o wymiarach 7,2×3,7 mm i rozstawie wyprowadzeń 0,25 mm. Współpracują z taśmami FPC o grubości 0,2 mm. Charakteryzują się napięciem i prądem znamionowym, odpowiednio, 30 V/0,2 A oraz zakresem dopuszczalnej temperatury pracy od -55 do +85°C.

www.hirose.com

Siłowniki z serii IHPT ze sprzężeniem haptycznym w wersjach o sile nominalnej do 120 N

Siłowniki ze sprzężeniem haptycznym z serii IHPT firmy Vishay są teraz dostępne w wersjach o sile nominalnej zwiększonej do 120 N. Oferta producenta obejmuje model IHPT1411AFELR73ABA

z kwalifikacją AEC-Q200, przeznaczony do zastosowań w motoryzacji (konsole, kierownice, dźwignie zmiany biegów, siedzenia) oraz cztery modele do aplikacji m.in. przemysłowych i medycznych (IHPT1207AGELR39AB0, IHPT1710ACELR2AB0, IHPT1411AFELR73AB0 i IHPT1614ACELR2R7BB0).



Nowe siłowniki z serii IHPT są polecane do zastosowań w środowiskach o wysokim poziomie hałasu, gdzie powiadomienia dźwiękowe mogą nie być słyszalne dla użytkownika. Pracują z nominalnym napięciem zasilania 12 V (dopuszczalny zakres 8...16 V), co – w odróżnieniu od innych podobnych mechanizmów ze sprzężeniem haptycznym – eliminuje konieczność stosowania dodatkowych układów wysokonapięciowych. Są w stanie nadać obciążeniu o masie 0,5 kg przyspieszenie do 6 g przy sterowaniu 12-woltowym impulsem elektrycznym o czasie trwania 5 ms.

Pozostałe parametry:

- indukcyjność cewki: od 1,8 mH,
- DCR: od 0,95 Ω,
- czas odpowiedzi: 5 ms,
- izolacja: 150 V_{DC},
- maksymalna temperatura robocza: +105°C,
- powierzchnia: od 29 × 21 mm do 44 × 37 mm.

www.vishay.com

Moduł komunikacyjny do aplikacji IoT wspierający standardy LTE-M/NB-IoT i DECT NR+

Nordic Semiconductor informuje o rozpoczęciu masowej produkcji zintegrowanego układu komunikacyjnego nRF9151 do aplikacji IoT. Jest to układ typu SiP (System-in-Package), mogący

REKLAMA



HAMMOND®

1550ZF odlewana obudowa aluminiowa z kołnierzem IP68

Dowiedz się więcej:
hammondmfg.com/1550zf



eusales@hammondmfg.com • +44 1256 812812





pełnić funkcję samodzielnego modemu komórkowego lub procesora aplikacyjnego i obsługującego standardy LTE-M/NB-IoT oraz DECT NR+. Jest obecnie najmniejszym i najbardziej energooszczędnym tego typu układem w ofercie firmy, mogącym znaleźć zastosowanie w automatyce przemysłowej, miernikach zużycia mediów, inteligentnym rolnictwie, aplikacjach smart city czy też w systemach śledzenia zasobów.

Moduł nRF9151 zawiera mikroprocesor aplikacyjny ARM Cortex-M33 64 MHz z 1 MB pamięci Flash i 256 kB pamięci RAM, modem LTE-M/NB-IoT, układ zarządzania zasilaniem i głowicę w.cz. Podobnie jak poprzednik o symbolu nRF9161, obsługuje standardy LTE-M/NB-IoT i DECT NR+, charakteryzuje się natomiast mniejszymi gabarytami, może pracować z większą mocą wyjściową w trybach Power class 3 (23 dBm) i Power class 5 (20 dBm) oraz oferuje mniejszy o 45% szczytowy pobór mocy. W przyszłych wydaniach oprogramowania firmware zostanie wprowadzona obsługa sieci NTN (Non-Terrestrial Network).

Moduł nRF9151 zapewnia kompatybilność z oprogramowaniem oraz istniejącymi narzędziami deweloperskimi opracowanymi dla nRF9160 i nRF9161. Jest w pełni wspierany przez platformę nRF Cloud firmy Nordic, oferującą usługi w chmurze dla produkowanych przez firmę modułów komunikacji bezprzewodowej, m.in. w zakresie lokalizacji, ochrony danych i zdalnego zarządzania.

www.nordicsemi.com

Tłumiki chipowe na pasmo 50 GHz w obudowach rozmiaru 0404 i 0604

Smiths Interconnect rozszerza ofertę chipowych tłumików wysokoczęstotliwościowych o nową serię TSX WB2, spełniającą wymogi normy niezawodnościowej MIL-PRF-55342. Elementy te mogą znaleźć zastosowanie w sektorze militarnym i lotniczym, są łatwe w implementacji i kompaktowe (dostępne wersje są produkowane w obudowach chipowych o rozmiarach 0404 i 0604). Pracują w zakresie częstotliwości do 50 GHz z maksymalną mocą sygnału wejściowego wynoszącą, w zależności od modelu, od 1 do 3 W. Zastosowana do ich produkcji technologia cienkowarstwowa na podłożach z azotku glinu zapewnia odporność na trudne warunki środowiskowe. Tłumiki z serii TSX WB2 wykazują mały współczynnik VSWR i wąski przedział tolerancji. Występują w wersjach o współczynniku tłumienia 1...10, 15, 20 i 30 dB.

www.smithsinterconnect.com

Mikroprzełączniki SMD SPST-NO o stopniu ochrony IP67

Mikroprzełączniki z serii EL2 zostały zaprojektowane do wymagających zastosowań z sektora przemysłowego, medycznego, automatyki budynkowej i elektroniki konsumenckiej. Są to elementy SMD

z wyprowadzeniami typu J-Bend, umożliwiające montaż za pomocą maszyn pick and place. Charakteryzują się stopniem ochrony IP67 i żywotnością 200 tys. cykli. Występują w wersjach o wysokości 3,5 mm i 5,2 mm oraz o sile nacisku 2,0 N i 3,5 N. Ich powierzchnia montażowa wynosi 6,2×6,2 mm. Zakres temperatury pracy rozciąga się od -40 do +85°C.

Pozostałe właściwości:

- obciążalność styków: 18 V_{DC}/150 mA,
- rezystancja izolacji: >100 MΩ,
- rezystancja styku: <100 mΩ,
- wytrzymałość dielektryczna: >250 V_{rms}.



www.littelfuse.com

Czujnik stężenia CO₂ do analizy jakości powietrza

Nowy czujnik stężenia CO₂, opracowany przez firmę Infineon, może znaleźć zastosowanie w analizatorach jakości powietrza wewnątrz pomieszczeń i umożliwiać poprawę sprawności energetycznej budynków. Został oparty na bazie technologii spektroskopii fotoakustycznej (PAS), zapewniającej wysoką dokładność pomiaru w czasie rzeczywistym, co pozwala m.in. na dynamiczne dostosowywanie parametrów pracy systemów wentylacyjnych do rzeczywistej liczby osób przebywających w pomieszczeniu, a w konsekwencji umożliwia obniżenie poboru mocy instalacji HVAC.



Technologia fotoakustyczna PAS bazuje na zjawisku polegającym na pochłanianiu przez cząsteczki gazu promieniowania podczerwonego, emitowanego przez wbudowane źródło światła. Proces ten generuje lokalne zmiany ciśnienia, mierzone przez czułe mikrofony, co pozwala na precyzyjne określenie stężenia CO₂. Zastosowanie technologii PAS pozwala na osiągnięcie wysokiej dokładności pomiaru przy jednoczesnym zachowaniu małych rozmiarów urządzenia.

XENSIV PAS CO2 5V jest zamykany w łatwej w integracji, miniaturowej obudowie o wymiarach 14×13,8×7,5 mm i pyłoszczelnej konstrukcji, zgodnej z wymogami normy ISO 20653:2013-02. Pracuje z napięciem zasilania 5 V, co dodatkowo ułatwia jego integrację w porównaniu z wcześniejszą wersją 12-woltową. Charakteryzuje się zakresem pomiarowym 0...32000 ppm i dokładnością ±50 ppm ±5% w podzakresie 400...3000 ppm. Skrócony do 55 s czas reakcji zapewnia szybsze działanie w środowiskach dynamicznych.

Czujnik może się komunikować z mikroprocesorem za pośrednictwem interfejsów UART, I²C lub PWM. Zastosowane w nim algorytmy kompensacji oraz funkcja automatycznej kalibracji ABOC (Automatic Baseline Offset Calibration) zapewniają dużą dokładność pomiaru, również w zmiennych warunkach środowiskowych. Zakres zastosowań XENSIV PAS CO2 5V obejmuje systemy HVAC, termostaty, inteligentne systemy oświetleniowe, zaawansowane urządzenia IoT (np. inteligentne lodówki) oraz inteligentne rolnictwo.

www.infineon.com

Pierwsze na rynku moduły kryptograficzne TPM z certyfikatem FIPS 140-3

STMicroelectronics wprowadza na rynek pierwsze moduły kryptograficzne STSAFE-TPM z certyfikatem FIPS 140-3, zastępującym wcześniejszą wersję FIPS 140-2. Wszystkie certyfikaty FIPS 140-2 mają wygasnąć we wrześniu 2026 r. Dzięki zgodności z FIPS 140-3, nowe moduły TPM pozwalają projektantom tworzyć bezpieczne, interoperacyjne urządzenia o wydłużonym czasie użytkowania



i certyfikacji. Obsługują funkcje bezpiecznego rozruchu (secure boot), zdalnej/anonimowej atestacji oraz bezpiecznej aktualizacji oprogramowania firmware, pozwalającej na dodawanie obsługi nowych algorytmów kryptograficznych, takich jak PQC. Dodatkowo zawierają większą wewnętrzną pamięć, przeznaczoną do przechowywania danych użytkownika (200 kB).

Układy ST33KTPM2X, ST33KTPM2XSPI, ST33KTPM2XI2C, ST33KTPM2I i ST33KTPM2A zapewniają kryptograficzną ochronę zasobów w krytycznych systemach informatycznych. Są przeznaczone do zastosowań w komputerach PC, serwerach i urządzeniach IoT komunikujących się z siecią, a także w aparaturze medycznej i innych aplikacjach o podwyższonych wymaganiach w zakresie bezpieczeństwa. ST33KTPM2I został zaprojektowany do systemów przemysłowych o długiej żywotności, natomiast ST33KTPM2A (ozn. STSAFE-V100-TPM) uzyskał kwalifikację AEC-Q100, pozwalającą na zastosowania w motoryzacji.

Nowe moduły STSAFE-TPM zapewniają kompatybilność z wieloma przemysłowymi standardami bezpieczeństwa, w tym Trusted

Computing Group TPM 2.0, Common Criteria EAL4+ i FIPS 140-3 level 1 z poziomem bezpieczeństwa fizycznego 3. Obsługują szyfrowanie w standardach ECDSA i ECDH (do 384 bitów), RSA (do 4096 bitów), AES (do 256 bitów) oraz SHA1, SHA2 i SHA3.

Firma STMicroelectronics oferuje usługi wstępnej konfiguracji, związane z wgrywaniem kluczy i certyfikatów. Pozwala to obniżyć koszty związane z integracją i zarządzaniem bezpieczeństwem, skraca czas wprowadzenia nowych produktów na rynek oraz zwiększa bezpieczeństwo łańcuchów dostaw.

www.st.com

Cewki TX i RX do systemów ładowania bezprzewodowego, odporne na wilgotność do 90% RH

Firma Vishay Dale zaprezentowała nowe cewki z rdzeniami ze sproszkowanego żelaza, przeznaczone do zastosowań w systemach ładowania bezprzewodowego o mocy do 30 W. Dzięki



nowo opracowanej powłoce ochronnej mogą one pracować przy wilgotności względnej do 90% RH. Ponadto wyróżniają się wysokimi wartościami dopuszczalnej temperatury pracy (do +105°C) i prądu nasycenia przy zachowaniu kompaktowych wymiarów oraz odpornością na wahania temperatury i nagłe spadki indukcyjności, obserwowane w odpowiednikach ferrytowych. W porównaniu zarówno z wcześniejszymi wersjami, jak i z odpowiednikami z oferty innych producentów zajmują mniejszą o około 25% powierzchnię na płycie drukowanej.

REKLAMA

INSTRUMENTY - TESTY I POMIARY

OSCYSKOPY | ZASILACZE | MULTIMETRY | TESTERY



KEYSIGHT
TECHNOLOGIES

GW INSTEK

pico
Technology

sensepeak

TT

**COMPUTER
CONTROLS**

Bielsko-Biała, ul. Bystrzańska 94

+48 33 485 94 90

info@ccontrols.pl
www.ccontrols.pl

	IWAS3222CZx	IWTX4646DCx	IWTX47R0DAx	IWTX47R0EBx
Indukcyjność	19,6 μ H	24 μ H	6,3 μ H	24 μ H
Moc znamionowa	5 W	30 W	30 W	30 W
Dobroć	28,5	185	190	200
Rezystancja DC	357 m Ω	75 m Ω	40 m Ω	75 m Ω
I_{RMS} (maks.)	1,2 A	7 A	7 A	6 A
I_{SAT} (maks.)	2,4 A	20 A	22 A	20 A
Wymiary	32x22 mm	46x46 mm	$\varnothing$ 47 mm	$\varnothing$ 47 mm
Grubość	3 mm	5,26 mm	5,3 mm	5,26 mm
Kształt	prostokątny	kwadratowy	okrągły	okrągły
Funkcja	Rx	Tx	Tx	Tx

Cewka odbiorcza IWAS3222CZEB190JR1 charakteryzuje się indukcyjnością 19,6 μ H, dobrocią 28,5 @ 200 kHz i rezystancją DC równą 357 m Ω . Cewki nadawcze IWTX4646DCB240JR1, IWTX47R0DAEB6R3JR1 i IWTX47R0EBEB240JR1 o indukcyjności z zakresu od 6,3 μ H do 24 μ H wykazują małą rezystancję DC, już od 40 m Ω , dobroć 185...200 i prąd znamionowy do 7 A. Wysoka wartość prądu nasycenia, sięgająca 22 A, zapewnia stabilną indukcyjność w szerokim zakresie natężenia prądu roboczego.

W celu zwiększenia trwałości do produkcji cewek nadawczych zastosowano przewód licowy pokryty jedwabiem, chroniącym przed zarysowaniami i okształceniami, zapewniającym lepszą izolację oraz zmniejszającym efekt naskórkowy. Ponadto modele IWTX4646DCB240JR1 i IWTX47R0EBEB240JR1 są produkowane przy użyciu techniki nawijania alfa, która eliminuje krzyżowanie się przewodów i pozwala uzyskać jak najmniejszy profil.

www.vishay.com

Akcelerometr cyfrowy MEMS do aplikacji przemysłowych o dużej dynamice

Firma TDK ogłasza wprowadzenie na rynek nowego akcelometru cyfrowego MEMS o symbolu Tronics AXO314, stanowiącego uzupełnienie platformy Tronics AXO300. Wcześniej do oferty firmy weszły modele: AXO315 – przeznaczony do awioniki, AXO305 (do lądowych i morskich systemów pozycjonowania) oraz AXO301 (do zastosowań w transporcie kolejowym i pomiarach nachylenia).

Nowy model, charakteryzujący się zakresem pomiarowym ± 14 g, został zaprojektowany z myślą o aplikacjach przemysłowych o dużej dynamice, wymagających odporności na wstrząsy i wibracje. Pracuje w zamkniętej pętli sprzężenia zwrotnego, co zapewnia bardzo dobrą liniowość i skuteczne tłumienie wibracji. Stanowi interesującą alternatywę dla droższych i większych akcelometrów kwarcowych klasy wojskowej.

W połączeniu z żyroskopem MEMS GYPRO4300, kompatybilnym pod względem rozkładu wyprowadzeń i interfejsu, AXO314 może znaleźć zastosowanie w wieloosiowych modułach IMU i INS klasy wojskowej, stosowanych m.in. w fotogrametrii i mapowaniu, a także w lotniczych i naziemnych systemach pozycjonowania ze wspomaganie GNSS. Jest kalibrowany fabrycznie w zakresie dopuszczalnej temperatury pracy od -40 do +85°C. Dane mogą być wyprowadzane również w postaci nieprzetworzonej (raw data), co pozwala na stosowanie dowolnych strategii kompensacji przez użytkownika.

Pozostałe właściwości:

- interfejs: 24-bitowy SPI,
- długoterminowa stabilność dryftu: 1 mg w skali roku,

- fluktuacje krótkoterminowe dryftu: 4 μ g,
- współczynnik tłumienia wibracji: 20 μ g/g²,
- gęstość szumu: 15 μ g/ $\sqrt{\text{Hz}}$,
- dryft w zakresie temperatury pracy: 0,5 mg,
- obudowa: ceramiczna SMD z 28 wyprowadzeniami J-LEAD.

www.tdk-electronics.tdk.com

Pierwszy satelitarne moduł komórkowy IoT-NTN z wbudowanym odbiornikiem GNSS

Sieci komórkowe pokrywają tylko niewielką część naszej planety, co powoduje rosnące zapotrzebowanie na rozwiązania zapewniające globalny, niezawodny zasięg w aplikacjach IoT do śledzenia zasobów w odległych lokalizacjach i na obszarach morskich. Satelitarne systemy IoT pomagają wypełnić tę lukę, jednak ich zastosowanie jest ograniczone przez wysokie koszty terminali. Firma u-blox wprowadza na rynek pierwszy satelitarne moduł IoT-NTN z wbudowanym odbiornikiem GNSS, oznaczony symbolem SARA-S528NM10. Został on oparty na chipsecie UBX-S52 i energooszczędnym odbiorniku GNSS M10. Spełnia wymagania 3GPP Release 17 dla globalnej łączności bezprzewodowej, zarówno w sieciach TN (terrestrial network), jak i NTN (non-terrestrial network). Dzięki zintegrowanemu odbiornikowi GNSS idealnie nadaje się do zastosowań wymagających ciągłego lub okresowego śledzenia zasobów, takich jak zarządzanie flotą pojazdów, inteligentne liczniki energii czy przemysłowe systemy sterowania.

Obecne systemy łączności satelitarnej zazwyczaj wymagają zastrzeżonego sprzętu i oprogramowania, co sprawia, że terminale są przypisane do określonych operatorów satelitarnych. Zmiana operatora wymagałaby również wymiany terminali. Rozwiązanie u-blox opiera się jednak na globalnych standardach 3GPP, co umożliwia współpracę z wieloma dostawcami usług satelitarnych. Nowy moduł SARA-S528NM10 został zaprojektowany z myślą o utrzymaniu łączności na obszarach pozbawionych zasięgu sieci komórkowych. Odbiornik GNSS zużywa mniej niż 15 mW mocy w trybie ciągłego śledzenia i oferuje wysoką czułość, co skraca czas potrzebny do ustalenia pozycji. Urządzenie dostarcza dane lokalizacyjne bez konieczności przerywania połączenia komórkowego lub satelitarnego, minimalizując tym samym pobór energii dzięki skróconemu czasowi aktywności urządzenia.

SARA-S528NM10 zapewnia rozszerzoną łączność przez LTE-M i NB-IoT w naziemnych sieciach komórkowych oraz NB-IoT w konstelacjach satelitów geostacjonarnych zgodnych z 3GPP Rel 17. Chipset UBX-S52 uzyskał certyfikację od Skylo, globalnego dostawcy usług NTN, do pracy w jego sieci satelitarnej. Moduł obsługuje trzy nowe pasma NTN: n23 (USA), n255 (globalne pasmo L) oraz n256 (europejskie pasmo S). Jest również kompatybilny z innymi modułami komórkowymi u-blox w formacie SARA, co ułatwia skalowanie istniejących produktów IoT bez konieczności czasochłonnego przeprojektowywania systemu.

www.u-blox.com

Miniaturowe przetłaczniaki ślizgowe do pracy w temperaturze otoczenia od -20 do +70°C

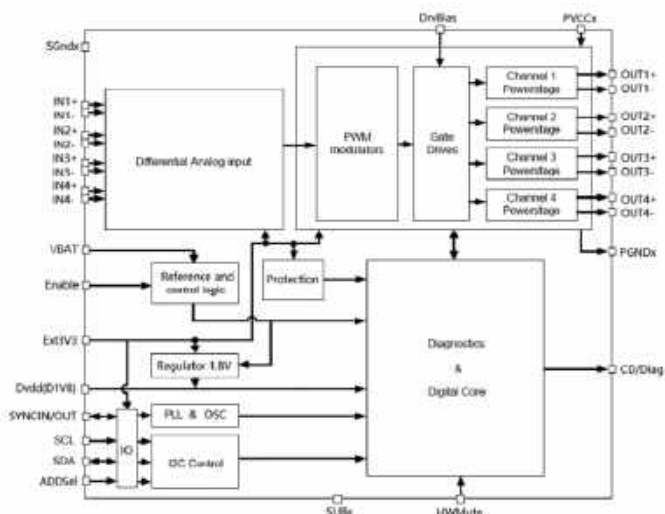
Firma Same Sky powiększa ofertę przetłaczniaków ślizgowych o ponad 60 nowych modeli, stanowiących rozszerzenie rodziny SLW. Są one dostępne w konfiguracjach SPDT, SP3T, SP4T, DPDT, DP3T i DP4T typu On-On, On-On-On, On-Off-On i Off-On-On-On. Dzięki małym gabarytom (8,3x5,6x3,6 mm) mogą znaleźć zastosowanie w urządzeniach o dużej gęstości upakowania



podzespołów, zarówno z sektora konsumenckiego, jak i przemysłowego. Opisywane komponenty są zgodne do pracy w temperaturze otoczenia od -20 do $+70^{\circ}\text{C}$.

Przełączniki z nowej serii występują w wersjach pionowych i poziomych o wysokości od 2,0 do 6,9 mm, przeznaczonych do montażu przewlekane. Charakteryzują się napięciem znamionowym 30 lub 50 V_{DC} i prądem znamionowym od 200 do 500 mA. Ich ceny hurtowe zaczynają się od 0,24 USD przy zamówieniach 500 sztuk.

www.sameskydevices.com



Wzmacniacz audio klasy D o mocy 4x49 W do elektroniki samochodowej

HFA80A to miniaturowy, 4-kanalowy wzmacniacz audio klasy D z kwalifikacją AEC-Q100, zaprojektowany do zastosowań w osprzęcie motoryzacyjnym. Układ jest produkowany w obudowie LQFP48L o powierzchni 7×7 mm i wymaga małej liczby elementów współpracujących. Zawiera wyjście mostkowe zrealizowane na tranzystorach MOSFET i pozwalające oddać do obciążenia maksymalną moc 4×49 W w przypadku współpracy z głośnikami o impedancji 2Ω i napięciu zasilającym 14,4 V.

Układ zapewnia skuteczne tłumienie tętnień na linii zasilającej (PSRR=80 dB) i wykazuje małe przesłuchy międzykanałowe (86 dB @ 1 W/1 kHz/4 Ω). Dzięki zastosowaniu techniki rozpraszania widma spread-spectrum charakteryzuje się ponadto niskim poziomem generowanych zaburzeń elektromagnetycznych (CISPR 25 level V) bez konieczności dołączania wyjściowych elementów filtrujących. Z kolei topologia feedback-before-filter i duża częstotliwość wyjściowego sygnału PWM (2 MHz) pozwala na stosowanie w ewentualnym filtrze wyjściowym elementów o małych gabarytach.

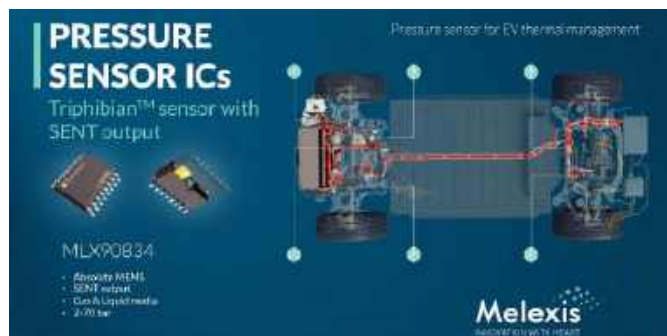
HFA80A wykazuje małą zawartość harmonicznych, już od 0,015% (@ 1 W, 1 kHz, 4 Ω) i małe znikowe napięcie szumu (typ. 52 μV przy wzmacnieniu 26 dB). Zapewnia płaską charakterystykę częstotliwościową do 40 kHz (lub 80 kHz po zoptymalizowaniu filtra). Dzięki małym opóźnieniom umożliwia realizację algorytmów tłumienia szumu.

Ceny hurtowe HFA80A zaczynają się od 4,80 USD przy zamówieniach 1000 sztuk.

Pozostałe cechy:

- napięcie zasilania: 4,5...18 V (odporność na przepięcia do 40 V),
- możliwość współpracy z głośnikami o impedancji 2Ω lub 4 Ω ,
- zabezpieczenie nadprądowe, przeciążeniowe, zwarciove i termiczne z 4 poziomami ostrzegawczymi,
- układ diagnostyczny wykrywający anomalie obciążenia,
- interfejs konfiguracyjny I²C,
- zabezpieczenie ESD do 2 kV HBM.

www.st.com

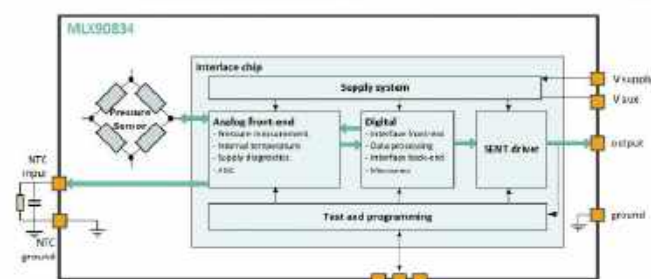


Czujnik ciśnienia cieczy i gazów z wyjściem SENT do zastosowań w motoryzacji

Melexis rozszerza rodzinę czujników ciśnienia Triphibian o nowy model MLX90834, zaprojektowany do zastosowań w motoryzacji, w tym w systemach HVAC, pompach i kompresorach oraz systemach monitorowania oleju silnikowego i przekładniowego. Jest to czujnik zrealizowany w technologii MEMS, umożliwiający pomiar ciśnienia cieczy i gazów (w tym czynników chłodniczych) w zakresie od 2 do 70 barów. Zawiera interfejs SENT umożliwiający odczyt wyników pomiaru ciśnienia bezwzględego i temperatury oraz danych diagnostycznych. Struktura układu obejmuje sensor mostkowy MEMS z elementami piezorezystancyjnymi (rozmieszczonymi na membranie), wejściowe obwody analogowe (front-end), mikrokontroler 16-bitowy i rdzeń DSP, realizujący m.in. kompensację temperatury, regulację napięcia i obsługę interfejsu SENT. Całość jest zamknięta w obudowie SO16 o wymiarach $10,2 \times 7,5 \times 2,3$ mm.

MLX90834 może współpracować z zewnętrznym termistorem NTC. Dzięki fabrycznej kalibracji zapewnia dokładność pomiaru temperatury lepszą niż $\pm 1^{\circ}\text{C}$. Dokładność pomiaru ciśnienia wynosi $\pm 0,5\%$ FS przez cały okres eksploatacji. Układ jest zabezpieczony przed przepięciami do 40 V i odwróceniem polaryzacji napięcia zasilającego. Może być stosowany w systemach o poziomie nienaruszalności bezpieczeństwa do ASIL C. Uzyskał kwalifikację AEC-Q100.

www.melexis.com



REKLAMA

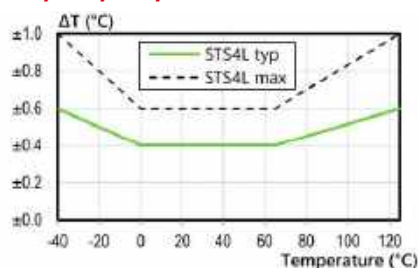
PRODUCENT
**ELEMENTÓW
INDUKCYJNYCH**

www.feryster.pl



Tani, cyfrowy czujnik temperatury o dokładności $\pm 0,4^{\circ}\text{C}$ i wymiarach $1,5 \times 1,5 \times 0,5$ mm

STS4L to najtańszy czujnik temperatury w ofercie firmy Sensirion, zaprojektowany z myślą o urządzeniach produkowanych masowo, w przypadku których ważne są koszty i niewielka przestrzeń montażowa. Sensor oferuje zakres pomiarowy od -40 do $+125^{\circ}\text{C}$ i dokładność $\pm 0,4^{\circ}\text{C}$ w zakresie $0...+60^{\circ}\text{C}$. Charakteryzuje się małym poborem mocy, pozwalającym na zastosowania w urządzeniach bateryjnych. Jest zamykany w obudowie SMD o wymiarach zaledwie $1,5 \times 1,5 \times 0,5$ mm.



STS4L komunikuje się przez interfejs I²C, taktowany zegarem do 1 MHz. Może być zasilany napięciem z zakresu od 1,08 do 3,6 V. Pobiera średnio 1,2 μA prądu przy pomiarach wykonywanych z interwałem 1 s.

Pozostałe parametry:

- pobór prądu w stanie spoczynkowym: 0,08 μA ,
- czas odpowiedzi (τ 63%): 2 s,
- dryft długoterminowy $< 0,03^{\circ}\text{C}$ na rok,
- odporność na wyładowania ESD: 2 kV (HBM).

www.sensirion.com



Stereofoniczny, 4-watowy wzmacniacz audio klasy D z 62-stopniową regulacją głośności

PAM8019E to stereofoniczny wzmacniacz audio klasy D z wyjściem mostkowym o mocy 4 W @ $R_L = 4 \Omega$ oraz z wbudowanym wzmacniaczem słuchawkowym 88 mW, pracującym w klasie AB. Może znaleźć

zastosowanie w monitorach LCD, odbiornikach TV, komputerach all-in-one i aktywnych głośnikach. Charakteryzuje się sprawnością sięgającą 90% oraz niskimi zniekształceniami i małym wyjściowym napięciem szumu ($\text{THD+N} = 0,03\%$, 40 μV rms przy maksymalnym wzmacnieniu). Jego wzmacnienie może być programowane (za pomocą sygnału napięciowego) w zakresie od -60 dB do $+20$ dB, w 62 krokach, co minimalizuje liczbę elementów współpracujących z układem. Z kolei modulacja spread spectrum pozwala ograniczyć poziom generowanych zaburzeń EMI, dzięki czemu do współpracy z układem można zastosować tanie filtry, bazujące na koralikach ferrytowych.

PAM8019E pracuje z napięciem zasilania od 2,5 do 6,0 V. Zawiera zabezpieczenie podnapięciowe, zwarcia i termiczne. Automatycznie wykrywa rozwarcie i zwarcie wyjścia podczas włączania zasilania, zapewniając ochronę całego toru audio jeszcze przed rozpoczęciem normalnej pracy. Technologia NCPL (non-clip power limit) automatycznie ogranicza moc wyjściową, poprawiając jakość dźwięku i zapewniając dodatkową ochronę głośników.

Układ jest produkowany w dwóch typach miniaturowych obudów: U-QFN3030-20 ($3,3 \times 3,3$ mm) i U-QFN4040-20 ($4,3 \times 4,3$ mm).

www.diodes.com



Nowe transceivery high-speed CAN FD SBC z 5-woltowym LDO w ofercie Microchip

Microchip Technology wprowadził do oferty nową rodzinę układów bazowych systemu ATA650x CAN FD (SBC) z w pełni zintegrowanym, szybkim transceiverem CAN FD i 5-woltowym regulatorem napięcia typu LDO. Komponenty są dostępne w kompaktowych obudowach 8-, 10- i 14-pinowych, o wymiarach odpowiednio: $2 \text{ mm} \times 3 \text{ mm}$ (VDFN8), $3 \text{ mm} \times 3 \text{ mm}$ (VDFN10) i $3 \text{ mm} \times 4,5 \text{ mm}$ (VDFN14).

Nowe transceivery CAN FD obsługują szybkość nadawania i odbioru danych do 5 Mb/s, a przy tym charakteryzują się bardzo niskim poborem mocy, przy typowym prądzie uśpienia wynoszącym zaledwie 15 μA . Układy ATA650x umożliwiają sterowanie napięciem zasilania za pomocą sygnałów magistrali, co zmniejsza pobór prądu przez jednostki sterujące (ECU) w nowoczesnych pojazdach. Aby jeszcze bardziej poprawić osiągi systemu w zakresie energooszczędności, nowe układy firmy Microchip są w stanie całkowicie odciąć zasilanie mikrokontrolera, wyłączając LDO, gdy moduł zostaje wprowadzony w tryb uśpienia.

Omawiane transceivery zostały wyposażone w zabezpieczenia, ESD oraz filtry EMC, a dodatkowo ułatwiają spełnienie wymogów w zakresie bezpieczeństwa funkcjonalnego zgodnie z normą ISO 26262 i schematem ASIL. Konstrukcja układów jest także zgodna z AEC-Q100 (ocena Grade 0). Zakres dopuszczalnych temperatur pracy ATA650x rozciąga się od -40°C do $+150^{\circ}\text{C}$.

www.microchip.com

dodaj do obserwowanych

Prezentujemy redakcyjny wybór najciekawszych projektów spośród ostatnio anonsowanych w internecie. Są to projekty na różnych etapach realizacji. Warto się zapoznać z projektami zakończonymi i śledzić realizację projektów niegotowych, by czerpać z nich inspirację do własnych prac.



Hobbystyczna maszyna pick & place

Prezentowany projekt to automatyczne urządzenie do precyzyjnego montażu elementów elektronicznych na płytkach drukowanych. Konstrukcja maszyny obejmuje system mechaniczny złożony z prowadnic liniowych, silników krokowych oraz głowicy podciśnieniowej, umożliwiającej chwyatanie komponentów i ich dokładne rozmieszczanie. Maszyna jest zintegrowana z oprogramowaniem, które kontroluje ruch i pozycję głowicy, automatyzując praktycznie cały proces układania podzespołów na powierzchni PCB.

Od strony software'u projekt został oparty na otwartoźródłowym projekcie OpenPnP. Sprzęt zbudowany przez autora wzorowany jest na kilku prezentowanych już w Internecie konstrukcjach, jednakże nie jest bezpośrednią ich kopią. Autor czerpie również inspirację z innych źródeł i dodaje pewne niestandardowe rozwiązania. Dodatkowo udało mu się uprościć niektóre komponenty i zwiększyć udział gotowych części, zmniejszając jednocześnie zapotrzebowanie na niestandardowe elementy, co powinno uczynić całość bardziej dostępną dla osób zainteresowanych odtworzeniem systemu we własnym zakresie.

<https://hackaday.io/project/175426-pick-and-place-machine>

Ultradokładny multimetr DIY

Ten imponujący projekt dotyka chyba wszystkich sekcji współczesnej elektroniki; system obejmuje bloki analogowe i układy programowalne FPGA, implementuje ponadto techniki cyfrowego przetwarzania sygnałów (DSP), a dodatkowo zawiera... tranzystory germanowe, często przestarzałe już półprzewodniki w obudowach metalowych, a także całkiem sporo zwykłych (i niezwykłych) elementów dyskretnych.

Urządzenie to multimetr stacjonarny. Teoretycznie sam miernik uniwersalny nie jest niczym skomplikowanym, ale jak to zwykle bywa w przypadku spotkania z rzeczywistością, prawda jest trochę inna. Ten konkretny multimetr to 6-cyfrowa aparatura, zdolna do pomiaru napięcia (w zakresach do 1, 10 i 100 V), prądu (w zakresach do 10 mA i 100 μ A, z oddzielnym zakresem 1 A) i rezystancji (1, 10, 100 i 1000 k Ω). Co ważne, jest to sprzęt 6-cyfrowy w dosłownym znaczeniu, gdyż mówimy tutaj nie tylko o liczbie prezentowanych na wyświetlaczu znaków, ale także o liczbie cyfr znaczących, czyli de facto o jego rzeczywistej rozdzielczości.



Główną częścią każdego współczesnego multimetru jest przetwornik analogowo-cyfrowy (ADC), który zamienia napięcie na wartość cyfrową, przetwarzaną następnie i wyświetlaną w odpowiedniej formie. Przetwornik ADC mierzy bezpośrednio napięcie, co realizuje funkcję woltomierza. Do pomiaru poza zakresem wejściowym przetwornika analogowo-cyfrowego wymagany jest dzielnik napięcia (umożliwiający pomiar wyższego napięcia) lub wzmacniacz (dla niższych napięć). W przypadku pomiaru prądu typową metodą jest zamiana natężenia na napięcie za pomocą rezystora szeregowego. Aby mieć możliwość zmiany zakresów pomiarowych, zwykle stosuje się układ kilku przełączanych rezystorów o różnej wartości.

Do pomiaru rezystancji autor zastosował dodatkowe źródło prądu, które przepuszcza prąd o znanym natężeniu przez mierzony element, a woltomierz dokonuje pomiaru spadku napięcia, który jest proporcjonalny do wartości rezystora. I to byłoby tyle, jeśli chodzi o prostszą część tego projektu – wszak większość multimetrów działa dokładnie w ten sam sposób. Jednakże, jak dowiadujemy się ze strony autora na portalu hackaday.io, znalazło się tutaj kilka przysłowiowych diabłów ukrytych w szczegółach.

Pierwszym z nich jest napięcie odniesienia ADC. Drugi stanowi izolacja galwaniczna – ponieważ mierzony układ może znajdować się na różnych potencjałach względem ziemi, klasyczne multimetry stacjonarne muszą być odizolowane od sieci. Trzecim jest oczywiście ADC, którego konstrukcja w zasadzie stanowi główny cel tego projektu. Na rynku istnieją stosunkowo dobre, zintegrowane przetworniki A/C (np. AD7177, LTC2380), ale autor tej konstrukcji zdecydował się na tradycyjną implementację dyskretnego ADC typu wielobocznego, stosowaną od dziesięcioleci w licznych miernikach o wysokiej rozdzielczości i nadal niedoścignioną przez scalone ADC pod względem wielu parametrów.

O tym, w jaki sposób wszystkie wymienione problemy w swojej konstrukcji rozwiązał autor opracowania, dowiedzieć się można z obszernego artykułu, pełnego logów z prototypowania różnych sekcji urządzenia. Obecnie autor nie uznaje jeszcze swojego projektu za w pełni zakończony. Jakkolwiek samo urządzenie już działa, jego twórca pracuje jeszcze nad porządną dokumentacją całości.

<https://hackaday.io/project/174022-diy-6-digit-multimeter>



W ofercie AVT*

AVT6067

Najważniejsze parametry:

- sumowanie sygnałów z dwóch źródeł sygnału stereofonicznego i jednego monofonicznego lub mikrofonu elektretowego,
- niezależna regulacja głośności każdego z wejść,
- wspólna regulacja głośności sumy kanałów (master volume),
- regulacja tonów niskich (bass) i wysokich (treble) sygnału wyjściowego (zsumowanego),
- dwa wyjścia: z regulacją barwy tonu i bez,
- płynna regulacja wzmocnienia tonów niskich oraz wysokich, z możliwością ich uwypuklenia oraz stłumienia,
- prosta budowa, łatwo dostępne elementy, zwarta konstrukcja,
- zasilanie: 12...24 V (DC),
- pobór prądu: ok. 20 mA.

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

- wersja [C] – zmontowany, uruchomiony i przetestowany zestaw [B] (elementy wstawiane w płytkę PCB),
 - wersja [A] – płytka drukowana bez elementów i dokumentacji.
- Kity, w których występuje układ scalony wymagający zaprogramowania, mają następujące dodatkowe wersje:
- wersja [A+1] – płytka drukowana [A] + zaprogramowany układ [UK] i dokumentacja,
 - wersja [UK] – zaprogramowany układ.

Dodatkowe materiały do pobrania ze strony www.ulubionykiosk.pl/media

- AVT5957 Sumator dwóch źródeł audio (EP 10/2022)
- AVT5873 Stereofoniczny aktywny regulator głośności (EP 8/2021)
- AVT5683 Trzykanałowy sumator/mikser audio (EP 6/2019)
- AVT1972 Potencjometr „Panorama” audio (EP 9/2017)
- AVT1958 Ducker audio z układem THAT4301 (EP 8/2017)
- AVT1670 Stereofoniczny regulator barwy dźwięku (EP 4/2012)
- AVT5208 T-Mixer. Nowoczesny mikser audio z panelem dotykowym (EP 11/2009)
- AVT2710 Prosty dyskotekowy mikser (EdW 2/2004)
- AVT490 Mikser audio ze sterowaniem cyfrowym (EP 2–3/1999)
- AVT2173 Trzykanałowy mikser ze wzmacniaczem (EdW 12/1997–1/1998)
- AVT1034 Czterokanałowy mikser stereo (EP 4/1995)
- AVT2132 Przedwzmacniacz z regulacją barwy dźwięku

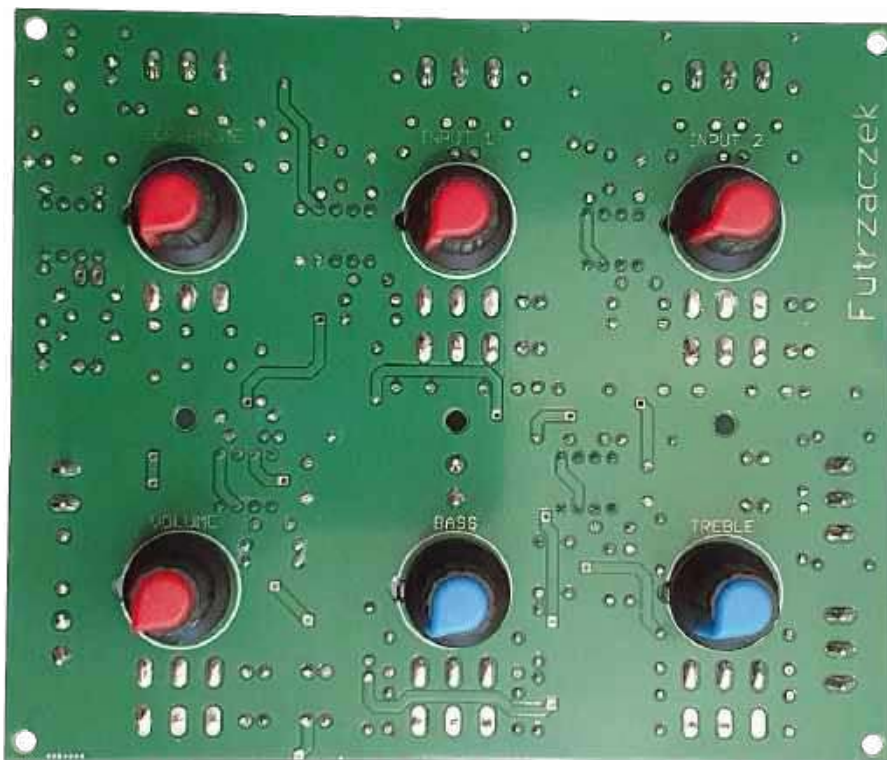
Nie każdy zestaw AVT występuje we wszystkich wersjach! Każda wersja ma załączony ten sam plik PDF! Podczas składania zamówienia upewnij się, którą wersję zamawiasz!
<http://sklep.avt.pl>

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

Mikser dyskotekowy trzech źródeł sygnału (1)

*Lewa strona, jak się bawicie?!
Prawa strona, jak się bawicie?!
Wszyscy ręce w górę! Aby poprowadzić imprezę, przyda się nie tylko charyzma, lecz również odpowiedni sprzęt. Nie zawsze drogi i wyposażony w mnóstwo funkcji – czasem amatorskie rozwiązania okazują się całkowicie wystarczające. W artykule opisujemy prosty mikser trzech źródeł sygnału audio, który może mieć wiele zastosowań, nie tylko rozrywkowe.*

Zaprezentowane w tym artykule urządzenie zostało nazwane mikserem dyskotekowym, lecz jego zastosowanie nie musi ograniczać się wyłącznie do tego obszaru. Amatorskie studio nagrań czy odtwarzanie muzyki i komunikatów przez mikrofon w lokalu handlowym to tylko przykładowe zastosowania. Jego obsługa jest tak prosta, że może służyć nawet dzieciom w rozwijaniu ich zainteresowań muzycznych! Nie trzeba kupować urządzenia bardzo rozbudowanego, wyposażonego w szeroki wachlarz funkcji, których początkujący nigdy nawet nie odkryje, nie mówiąc już o poprawnym użyciu.



To wszystko dzięki garści popularnych elementów, które można kupić w dosłownie każdym sklepie elektronicznym. Zwarta konstrukcja ułatwia montaż i niweluje konieczność oddzielnego lutowania

każdego z kilkudziesięciu cienkich przewodów do potencjometrów. Zasilanie i złącza sygnałowe – to wszystko, czego od świata zewnętrznego potrzebujemy, czego od światła zewnętrznego potrzebujemy, czego od światła zewnętrznego potrzebujemy, czego od światła zewnętrznego potrzebujemy.

Wykaz elementów:

Rezystory: (THT o mocy 0,6 W, 1%)
R1, R2, R49: 2,2 kΩ
R3, R4, R6...R32, R41, R42, R47, R48: 100 kΩ
R5: 4,7 kΩ
R33...R37, R39, R43, R44: 10 kΩ
R38, R40: 3,3 kΩ
R45, R46: 47 Ω
P1: 10 kΩ jednoobrotowy, pojedynczy, logarytmiczny, na panel
P2...P4: 10 kΩ jednoobrotowy, podwójny,

logarytmiczny, na panel
P5, P6: 100 kΩ jednoobrotowy, podwójny, liniowy, na panel

Kondensatory:
C1, C3, C7...C10, C15...C18, C23, C24, C44...C49: 470 nF (63 V, R=5 mm, MKT)
C2, C5, C6, C13, C14, C21, C22, C27...C30, C41...C43: 100 µF (25 V, R=2,5 mm)
C4: 150 pF (monolityczny, R=5 mm)

C11, C12, C19, C20, C25, C26, C39, C40: 33 pF (monolityczny, R=5 mm)
C31...C34: 33 µF/63 V (R=5 mm MKT)
C35...C38: 3,3 µF/63 V (R=5 mm MKT)
C50, C51: 470 µF/35 V (R=5 mm)

Półprzewodniki:

D1: 1N5819
U1...U5: TL082 (DIP8, opis w tekście)

Pozostałe:

J1...J5: złącze śrubowe ARK3/500
J6: złącze śrubowe ARK2/500
JP1: goldpin męski 2 pin 2,54 mm THT pionowy + zworka
5 szt. podstawek DIP8
2 szt. niebieskich galek do potencjometrów GAŁ-7877 (opis w tekście)
4 szt. czerwonych galek do potencjometrów GAŁ-12477 (opis w tekście)

Budowa

Schemat ideowy omawianego układu znajduje się na **rysunku 1**. Pierwszym blokiem, jaki można na nim wyróżnić, jest wejście monofoniczne, mogące obsługiwać sygnał z mikrofonu. Zasilanie pojedynczego przetwornika może pochodzić z trzeciego zacisku złącza J1, na który trafia solidnie odfiltrowane napięcie stałe (przez filtr dolnoprzepustowy na rezystorze R1 i kondensatorach C1 i C2), przechodzące przez rezystor R2, na którym może odkładać się napięcie wytworzone przez wkładkę mikrofonową. Wzmacniacz US1A zapewnia wzmocnienie około 21 V/V, co odpowiada wartości nieco ponad 26 dB w skali logarytmicznej. Kondensator C3 odcina składową stałą na wejściu, zaś rezystor R3 ustala ją na nowo, na wartość równą połowie napięcia zasilającego. Kondensator C4 zawęża pasmo przenoszenia wzmacniacza do nieco ponad 10 kHz, by niepotrzebnie nie przenosił szumu generowanego przez mikrofon. Jeżeli układ nie musi cechować się wzmocnieniem, ponieważ np. zastosowany mikrofon lub inne źródło sygnału ma już wbudowany przedwzmacniacz napięciowy, można zewrzeć zwórkę JP1 i ograniczyć wzmocnienie tego bloku do 1 V/V, czyli 0 dB. Potencjometrem P1 można regulować udział tego wejścia w zsumowanym sygnale wyjściowym.

Następne dwa bloki, które obsługują wejścia stereofoniczne, są takie same. Rezystory 100 kΩ obciążają wyjścia źródeł sygnału dla składowej stałej, choć nieznacznie, gdyż ich rezystancja jest relatywnie wysoka – jest to jednak konieczne dla niektórych typów wyjść. W ustalaniu impedancji wejściowej biorą również udział rezystory wejściowe wzmacniaczy odwracających. Mamy w ten sposób kontrolę nad impedancją wejściową, która wynosi 50 kΩ, ale również odciętą ewentualną składową stałą, pochodzącą ze źródeł sygnału. Wzmocnienie buforów wejściowych wynosi –1 V/V, zaś ich główną funkcją jest ustalenie impedancji zasilającej potencjometry. Gdyby nie one, wpływ na głośność danego wejścia miałaby impedancja wyjściowa źródła sygnału w stopniu znacznie większym, niż ma to miejsce teraz. Kondensatory 33 pF ograniczają pasmo przenoszenia układu, ale znacznie słabiej, niż w przypadku wejścia mikrofonowego – obliczone teoretycznie pasmo trzydecybelowe wynosi około 48 kHz. Nie ma to zatem wpływu na pasmo akustyczne, za to zmniejsza słyszalny poziom szumu.

Wyjścia wzmacniaczy operacyjnych we wszystkich blokach są obciążone rezystorami o wartości 100 kΩ do masy w celu zlinearyzowania stopni wyjściowych wzmacniaczy operacyjnych. Ponieważ pracują one przy zasilaniu asymetrycznym, potencjał

sztucznej masy jest równy połowie napięcia zasilającego, więc takie samo napięcie będzie panować między wyjściem każdego ze wzmacniaczy operacyjnych a masą. Doświadczenie pokazuje, że obciążenie wyjścia dla składowej stałej poprawia odpowiedź impulsową układu, ponieważ przez tranzystor stopnia wyjściowego ciągle płynie składowa stała prądu.

Do regulacji głośności użyto potencjometrów jednoobrotowych 10 kΩ, w których składowa stała prądu wyjściowego wzmacniaczy operacyjnych została obciążona kondensatorami elektrolitycznymi 100 μF, tworząc filtr górnoprzepustowy o częstotliwości granicznej około 0,15 Hz, co jest w zupełności wystarczające do tego, by w użytecznym paśmie pracy układu (zaczynającym się od około 20 Hz) nie było zauważalnego przesunięcia fazy. Dzięki temu, że przez ścieżkę oporową potencjometru nie płynie prąd stały, nie wprowadza on do sygnału trzasków wynikających z przerywania drogi prądu przez ślizgacz trący po niejednorodnej powierzchni ścieżki oporowej.

Filtry górnoprzepustowe, odcinające składową stałą, znajdują się również na wejściach układu, zaś ich częstotliwość odcięcia można określić na nie więcej niż 3,4 Hz (im wyższa impedancja wewnętrzna źródła, tym ta częstotliwość będzie niższa). Z uwagi na pracę z wysokimi rezystancjami użyto w tych miejscach kondensatorów z dielektrykiem foliowym. Gdyby zastosować kondensatory elektrolityczne, wówczas ich prąd upływu mógłby odbić się negatywnie na separacji składowej stałej i cały ten wysiłek poszedłby wówczas na marne.

Po unormowaniu impedancji wyjściowej każdego ze źródeł sygnału i/lub jego wzmocnieniu przez przedwzmacniacz mikrofonowy, sygnały są dzielone przez potencjometry i trafiają do sumatora sygnałów. Aby nie przenikały one między sobą wzajemnie, użyto sumatora w postaci wzmacniacza odwracającego, ponieważ w punktach podłączonych do wejść odwracających układów US4A i US4B potencjały powinny się (teoretycznie) trzymać stałej wartości, niezależnie od przyłożonego sygnału, ponieważ wynika to z warunku pracy wzmacniacza operacyjnego objętego ujemnym sprzężeniem zwrotnym. Zatem wszystkie te trzy sygnały będą spływały do punktów, które one będą traktowały jako wirtualną masę o niezmiennym potencjale. Ma to jeszcze sens

pod tym względem, że bufor wejściowy na wcześniejszym etapie torów przetwarzania sygnału mają charakter odwracający i tutaj również mamy wzmacniacz o charakterze odwracającym – faza sygnału między wejściem a wyjściem pozostanie zatem bez zmian.

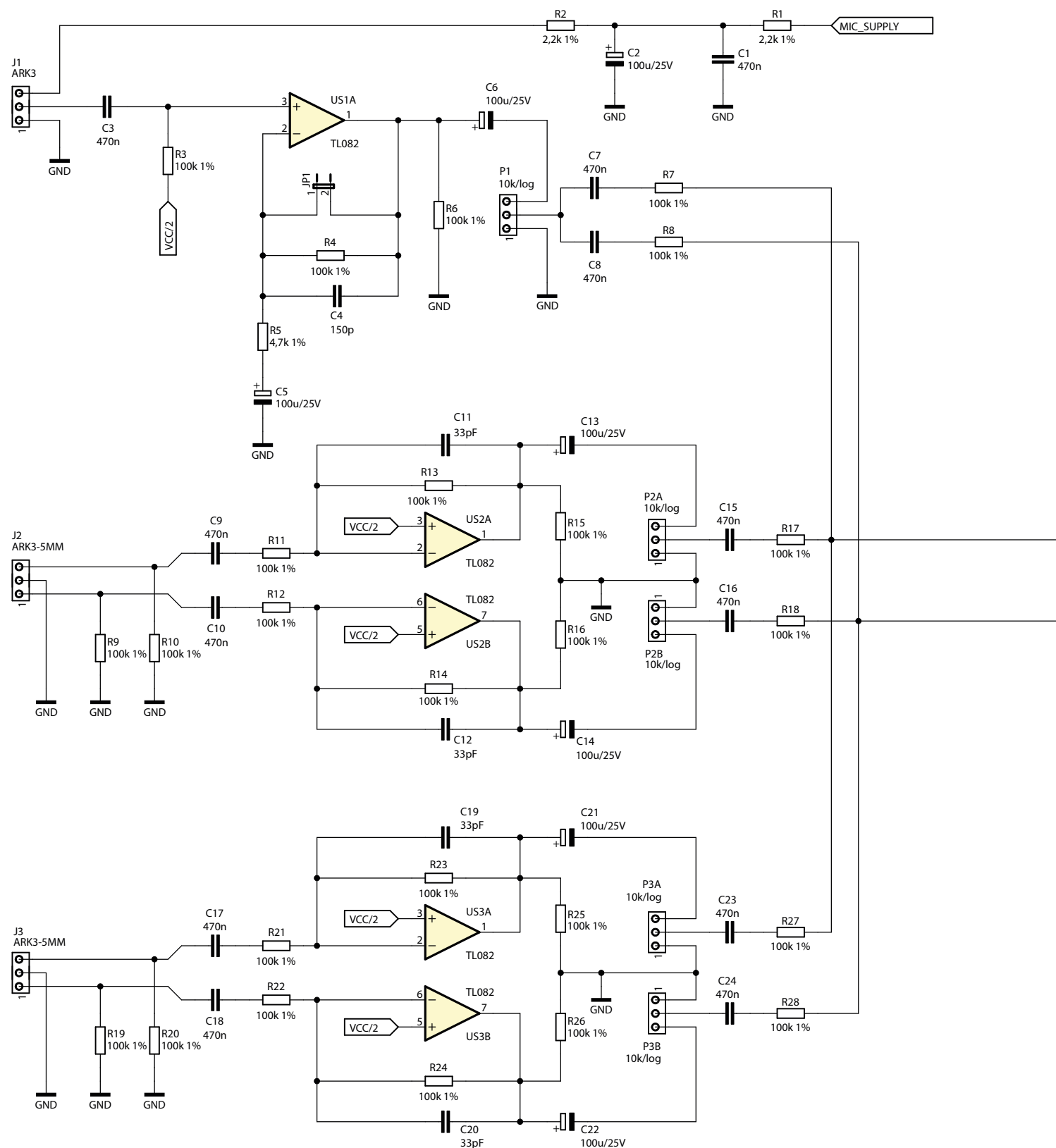
Zarówno na wejściach, jak i na wyjściach sumatora również konieczna była separacja składowej stałej. Mimo, że układ można byłoby zaprojektować w wersji z przenoszeniem składowej DC, to w praktyce okazuje się to bardzo trudne, ponieważ offset napięciowy wzmacniaczy operacyjnych powoduje powstawanie różnic rzędu kilku miliwoltów między poszczególnymi wyjściami. To z kolei wpływa na pracę następnych stopni, a ponadto wymaga takiego zaprojektowania układu, by skompensować wpływ prądów wejściowych wzmacniaczy operacyjnych – co w przypadku układów TL082 nie jest szczególnie istotne (wejścia z tranzystorami JFET), lecz przy wymianie tych układów na inne taki problem mógłby już dojść do głosu. Separacja składowej stałej na każdym kroku wprawdzie zwiększa liczbę elementów, czyni za to układ prostszym w uruchomieniu i łatwiejszym do modyfikacji.

Potencjometr P4 realizuje funkcję, którą dźwiękowcy określiliby mianem „master volume”, a która w rzeczywistości jest regulacją głośności sumy wszystkich trzech kanałów. Ten sygnał jest dostępny na zaciskach złącza J4 – impedancja wyjściowa jego źródła jest relatywnie niska, nie przekracza bowiem 5 kΩ.

Ostatnim członem jest prosty regulator barwy tonu, umożliwiający zarówno stłumienie, jak i podbicie tonów niskich i wysokich. Podwójne potencjometry zastosowane w tym miejscu pozwalają na uzyskanie współbieżnej regulacji w obu kanałach. Ewentualne nierównomierności między ich sekcjami nie będą drastycznie odczuwalne, bowiem ludzkie ucho jest



Fotografia 1. Widok zmontowanego układu od strony elementów



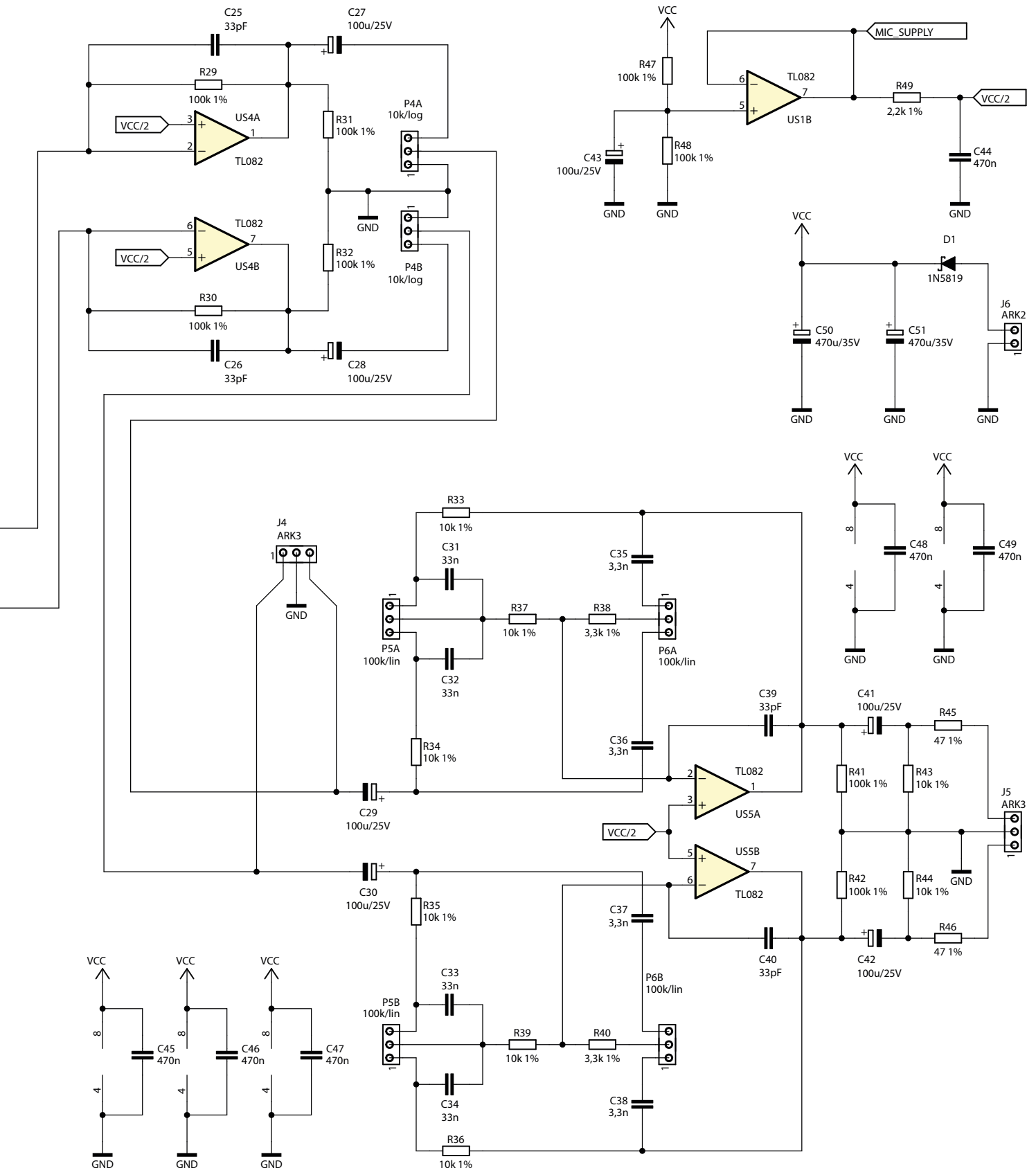
Rysunek 1. Schemat ideowy układu miksera dyskotekowego

słabo wyczułone na niewielkie rozbieżności w charakterystykach częstotliwościowych między lewym i prawym kanałem toru audio. Impedancja wyjściowa tego stopnia jest bliska zeru, więc można z niego sterować następnym urządzeniem bez konieczności stosowania dodatkowych wtórników. Kondensatory elektrolityczne C41 i C42

odcinają składową stałą, która była dotychczas utrzymywana dla prawidłowej pracy wzmacniaczy operacyjnych, zaś rezystory R43 i R44 odpowiadają za prawidłową polaryzację tychże kondensatorów. Warto zauważyć, że rezystancja R43 i R44 jest dziesięciokrotnie niższa niż typowo używanych rezystorów 100 kΩ w układzie,

a to ze względu na skrócenie czasu dochodzenia do stanu ustalonego zerowej składowej stałej napięcia na wyjściu.

Zadaniem R45 i R46 jest dopasowanie impedancji wyjściowej układu do impedancji charakterystycznej kabla ekranowanego łączącego „klocki” audio. Unikamy w ten sposób również ryzyka



(trudnego do okiełznania) wzbudzenia się wzmacniaczy operacyjnych przy obciążeniu ich wyjść znaczącą pojemnością długiego przewodu. Zatem złącze J5 stanowi drugie wyjście układu, tym razem z modyfikowalną charakterystyką częstotliwościową. W jakim stopniu, pokażą to pomiary przeprowadzone na prototypie, o czym dalej.

W układzie znajduje się również prosty obwód wytwarzający „sztuczną masę”, czyli wtórnik napięciowy na dotychczas niewykorzystanej połowie układu US1, który odwzorowuje napięcie pochodzące z dzielnika napięciowego R47+R48. Owe napięcie jest filtrowane przy użyciu C43, który zapewnia wysoką stałą czasową filtracji.

Na potrzeby zredukowania wartości skutecznej szumu generowanego przez sam wzmacniacz operacyjny US1B, za jego wyjściem znalazł się człon całkujący R49+C44, obniżający znacznie poziom szumów przy jednoczesnym ustaleniu rezystancji wyjściowej tego obwodu na 2,2 kΩ.

Michał Kurzela, EP



W ofercie AVT*

AVT6068

Najważniejsze parametry:

- niskoszumowy, dwustopniowy przedwzmacniacz stereo,
- charakterystyka RIAA,
- pojemność wejściowa ustawiana w zakresie 47...200 pF,
- rezystancja wejściowa ustawiana przełącznikiem (47 kΩ lub 100 kΩ),
- regulacja wzmacnienia: 24/27/30 dB,
- zasilanie: ±15 V (DC)/200 mA (patrz opis).

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

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

Dodatkowe materiały do pobrania ze strony www.ulubionykiosk.pl/media

- AVT6062 Filtr częstotliwości podakustycznych „subsonic” (EP 11/2024)
- Wzmacniacz transkonduktancyjny do subwoofera (EP 7/2023)
- AVT5827 Przedwzmacniacz lampowy do gramofonu (EP 12/2020)
- AVT5634 Lampowy przedwzmacniacz gramofonowy (EP 8/2018)
- AVT5514 DSP1701_SUB cyfrowy filtr do subwoofera aktywnego (EP 10/2015)
- AVT1693 Przedwzmacniacz gramofonowy MM RIAA (EP 8/2012)
- AVT1687 Filtr do subwoofera (EP 8/2012)

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

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

Przedwzmacniacz gramofonowy MM z pasywną korekcją

Winylomania ma się dobrze, historia zatoczyła koło i coraz częściej w domowych systemach audio gramofon zastępuje odtwarzacz CD. Niestety współczesne wzmacniacze sporadycznie wyposażone są w wysokiej jakości wejście współpracujące z wkładką gramofonową. Prezentowany układ pełni funkcję przedwzmacniacza do wkładki z ruchomym magnesem (MM) i współpracuje z wejściem liniowym każdego wzmacniacza audio. Możliwość wyboru pojemności wejściowej oraz wzmacnienia zapewnia optymalną współpracę z wkładkami gramofonowymi MM różnych producentów.

Schemat przedwzmacniacza pokazano na **rysunku 1**. Układ oparty jest o wysokiej jakości wzmacniacza operacyjnego typu AD8599 firmy Analog Devices, zoptymalizowanym pod kątem niskich zniekształceń oraz znikomego poziomu szumów.

Schematy kanałów lewego i prawego są identyczne, elementy kanału lewego oznaczone są sufiksem *L, zaś prawego: *R. Opis działania dotyczy kanału lewego. Sygnał wejściowy z wkładki MM doprowadzony jest do złącza INL, a koralek ferrytowy FB1L filtruje zakłócenia w.c.z. Wstępne obciążenie pojemnościowe 47 pF, wymagane do poprawnej pracy wkładki, zapewnia kondensator C1L. Pojemność może zostać zwiększona zgodnie z wymaganiami zawartymi w dokumentacji używanego przetwornika, za pomocą sekcji 1,2,3 przełącznika DIP oznaczonego SW1L, poprzez równoległe podłączenie z C1L dodatkowych pojemności C2L...C5R.

Sumaryczna pojemność obciążająca może być regulowana od 47 pF do ok. 400 pF, z krokiem 47 pF. Podczas dopasowywania pojemności wejściowej każdorazowo należy pamiętać o uwzględnieniu pojemności przewodu połączeniowego, która może wynosić od kilkunastu do kilkuset pF. Obciążenie rezystancyjne wkładki zapewnia rezystor R1L oraz – odłączany sekcją 4 przełącznika SW1L – opcjonalny rezystor R2L. Umożliwia to zapewnienie typowej impedancji obciążenia 47 kΩ (SW1L-4 zwarte, położenie domyślne) oraz 100 kΩ.

Z wyjścia sekcji dopasowania obciążenia i separacji składowej stałej przez C6L, sygnał doprowadzony jest do pierwszego stopnia wzmacniacza o ustalonym przez rezystory R4L, R5L wzmacnieniu ok. 30 dB. Wartości rezystorów dobrane są jako kompromis pomiędzy obciążeniem wyjścia wzmacniacza operacyjnego U1L a możliwie minimalnym poziomem szumów. Kondensator



CE3L redukuje wzmacnienie stałoprądowe oraz zapewnia wstępne odcięcie sygnałów o bardzo niskich częstotliwościach, pełniąc funkcję podstawowego filtra subsonicznego/antywibracyjnego. W opisanym przedwzmacniaczu zastosowano korekcję pasywną – elementy kształtujące charakterystykę przenoszenia nie są umieszczone tradycyjnie w pętli sprzężenia zwrotnego, ale pomiędzy stopniami wzmacniacza. Każde z wymienionych rozwiązań ma pewne wady i zalety oraz – jak to zwykle bywa – ma swoich zagorzałych zwolenników i przeciwników. Z inżynierskiego punktu widzenia największą wadą układu pasywnego jest zwiększona komplikacja układu, gdyż musi on być zrealizowany na drodze dwóch stopni wzmacnienia, co wiąże się też z większymi zniekształceniami i poziomem szumów, zaletą zaś

Wykaz elementów:

Rezystory: (SMD MELF, 1%, metalizowane)
R3L, R3R, R4L, R4R, R10L, R10R, R14L, R14R: 220 Ω
R11L, R11R: 3,3 kΩ
R12L, R12R: 1,5 kΩ
R13L, R13R: 2,0 kΩ
R15L, R15R: 470 kΩ
R1L, R1R: 100 kΩ
R2L, R2R: 88,7 kΩ
R5L, R5R: 6,8 kΩ
R6L, R6R: 12 kΩ
R7L, R7R: 100 Ω
R8L, R8R: 1,2 kΩ

R9L, R9R: 560 Ω

Kondensatory:

C1L, C1R, C2L, C2R: 47 pF (foliowy, R=5 mm, 5% *)
CE1, CE2: 47 μF/25 V (elektrolityczny, low ESR, fi=5 mm, R=2 mm)
CE1L, CE1R, CE2L, CE2R: 4,7 μF/25 V (SMD 3216, tantalowy)
CE3L, CE3R: 220 μF/25 V (elektrolityczny, low ESR, fi=8 mm, R=3,5 mm)
C3L, C3R, C4L, C4R, C5L, C5R: 100 pF (foliowy, R=5 mm, 5% *)

CE4L, CE4R: 100 μF/25 V (elektrolityczny, low ESR, fi=6,3 mm, R=2,5 mm)
C6L, C6R: 4,7 μF (foliowy, R=5 mm, 5% *)
C8L, C8R, C9L, C9R: 0,1 μF (foliowy, R=5 mm, 5% *)
C10L, C10R: 15 nF (foliowy, R=5 mm, 5% *)
C11L, C11R, C13L, C13R: 47 nF (foliowy, R=5 mm, 5% *)
C12L, C12R: 33 nF (foliowy, R=5 mm, 5% *)
C14L, C14R: 100 nF (SMD 0805, 25 V, X7R)

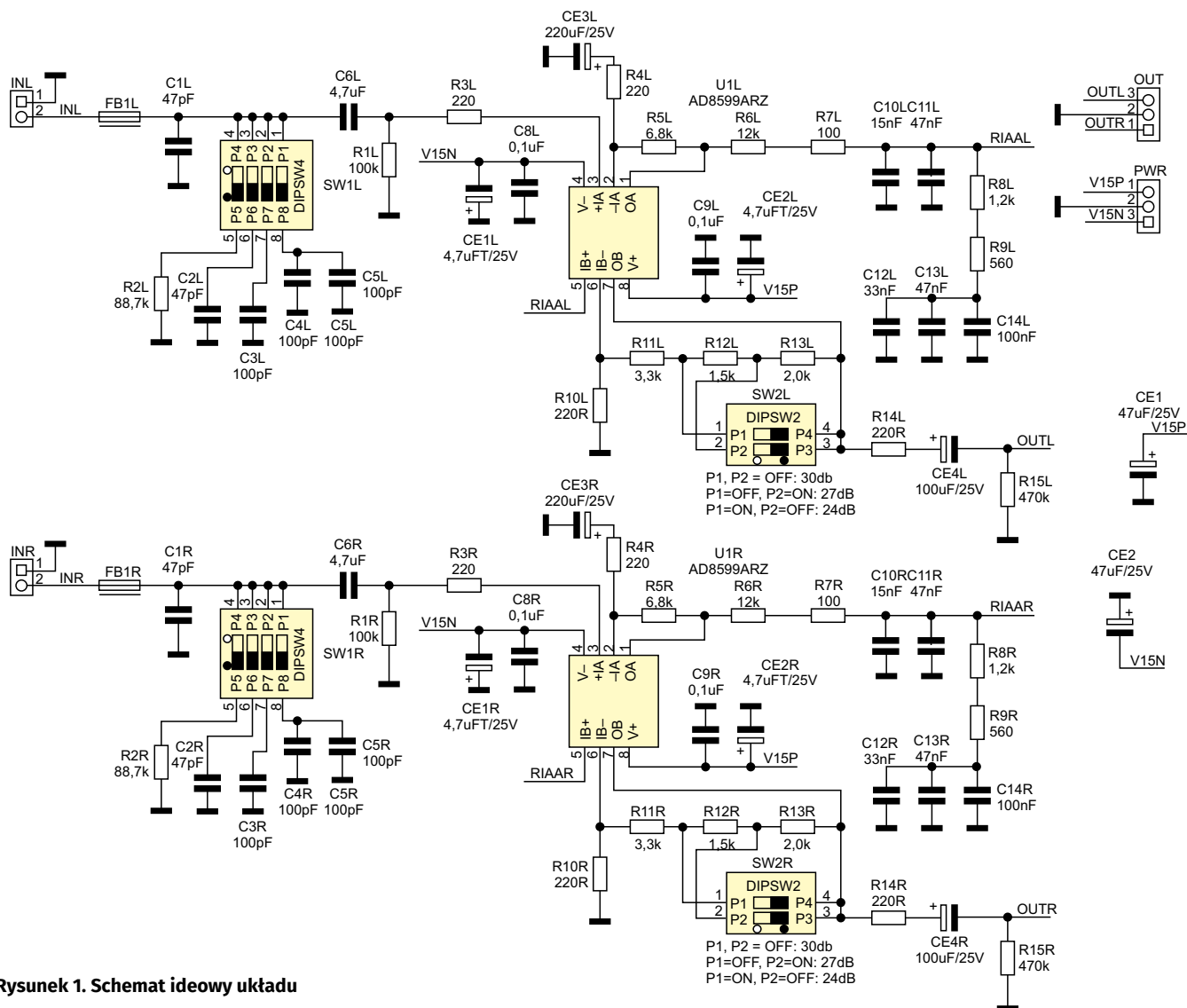
Półprzewodniki:

U1L, U1R: AD8599ARZ (SO8)

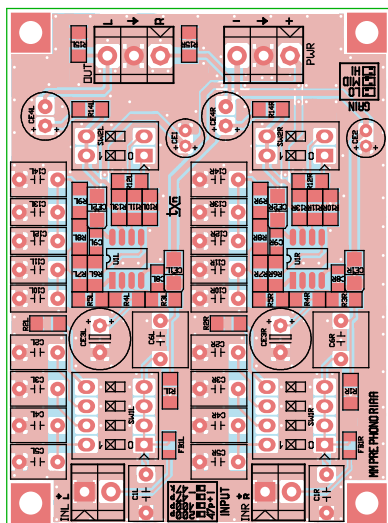
Pozostałe:

FB1L, FB1R: dławik ferrytowy BL-M21AH6015N1D (SMD 0805)
INL, INR: złącze DG 3,5 mm 2 pin (DG381-3.5-2)
OUT, PWR: złącze DG 3,5 mm 3 pin (DG381-3.5-3)
SW1L, SW1R: przełącznik DIP (4 sekcje)
SW2L, SW2R: przełącznik DIP (2 sekcje)

* dobrać – patrz opis



Rysunek 1. Schemat ideowy układu



Rysunek 2. Rozmieszczenie elementów na PCB

jest teoretyczny brak zniekształceń związanych z pracą korekcji w obwodzie sprzężenia zwrotnego. W modelu za korekcję RIAA odpowiadają elementy R6L...R9L oraz C10L...C13L. Równoległe i szeregowo połączenia elementów w obwodzie

korekcji zastosowano dla ułatwienia doboru elementów z typowego szeregu E12, co jest szczególnie istotne w przypadku kondensatorów.

Po pasywnej korekcji charakterystyki sygnału, ze względu na obniżenie poziomu, doprowadzony jest do drugiego stopnia U1LB o skokowo regulowanym wzmacnieniu, pozwalającym dostosować całkowite wzmacnienie toru do wkładek o różnych poziomach wyjściowych. Wzmacnienie regulowane jest poprzez wybór odpowiednich wartości rezystorów R11L...R13L w torze sprzężenia zwrotnego za pomocą przełącznika SW2L. Dostępne wartości wzmacnienia to ok. 24/27/30 dB. Sygnał z drugiego stopnia – po separacji składowej stałej (CE4L) – doprowadzony jest do złącza wyjściowego OUT. Układ może być zasilany symetrycznym napięciem ± 15 V doprowadzonym do złącza PWR.

Przedwzmacniacz zmontowano na dwustronnej płytce drukowanej, rozmieszczenie elementów pokazuje **rysunek 2**.

Sposób montażu jest klasyczny i nie wymaga opisu. Należy zadbać tylko

o zastosowanie niskoszumowych rezystorów metalizowanych o tolerancji 1% oraz dokładny dobór kondensatorów foliowych odpowiadających za kształtowanie charakterystyki. Poprawnie zmontowany przedwzmacniacz nie wymaga specjalnego uruchamiania. Ze względu na niski poziom sygnałów w torze należy zastosować ekranowane przewody połączeniowe, a płytkę umieścić w ekranowanej obudowie. W celu osiągnięcia optymalnych parametrów układ powinien być zasilany za pomocą niskoszumowego zasilacza liniowego ± 15 V o wydajności 200 mA.

Zmontowaną płytkę przedwzmacniacza można zobaczyć na **fotografii tytułowej**.

Przedwzmacniacz MM, uzupełniony o opisaną na łamach EP8/19 płytkę przedwzmacniacza gramofonowego dostosowanego do wkładek MC, niskoszumowego zasilacza audio (EP 10/2018) oraz filtra subsonicznego (EP 11/2024) może być podstawą do budowy uniwersalnego przedwzmacniacza gramofonowego DIY wysokiej jakości.

Milego odsłuchu!

Adam Tatuś, EP



W ofercie AVT*

AVT6069

Najważniejsze parametry:

- moc wyjściowa: 2...4 W (głośnik 4 Ω), 2...3 W (głośnik 8 Ω),
- pasmo przenoszenia: 40 Hz...20 kHz,
- poziom zniekształceń harmoniczných: 5%,
- napięcie zasilania: 9...18 V (patrz tekst).

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

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

Dodatkowe materiały do pobrania ze strony www.ulubionykiosk.pl/media

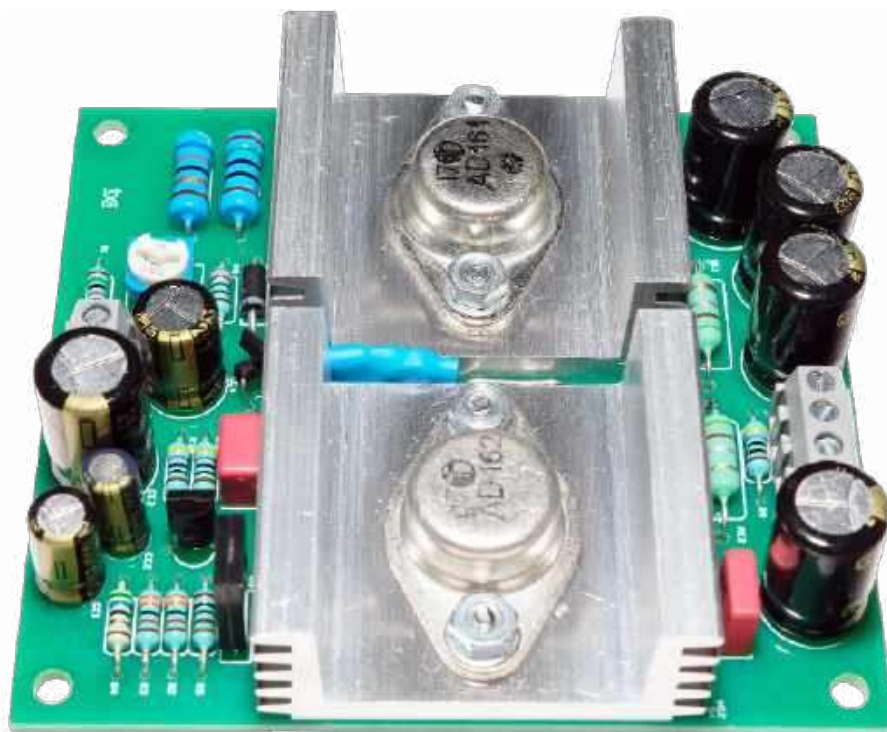
Projekt 260	Wzmacniacz lampowy PCL86 Stereo (EP 4/2024)
AVT6001	Sansuix – lampowy wzmacniacz mocy 2x20 W (EP 8–9/2023)
AVT5827	Przedwzmacniacz lampowy do gramofonu (EP 12/2020)
----	Automatyczny wyciszacz dźwięku po zaniku zasilania (EP 9/2020)
----	Wzmacniacz lampowy z regulacją barwy dźwięku (EP 5/2020)
Projekt 248	Wzmacniacz lampowy CFB (cathode feedback) (EP 11/2019)
AVT5727	Hybrydowy wzmacniacz słuchawkowy na lampie Nutube 6P1 (EP 11/2019)

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

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

Retro wzmacniacz mocy z tranzystorami germanowymi

Układy retro wykonane z zastosowaniem elementów z początków ery elektroniki cieszą się niesłabnącą popularnością. Przykładami mogą być kopie lub własne opracowania wzmacniaczy bazujących na lampach elektronowych. Wiele osób obawia się pracy z wysokimi napięciami o wartości setek woltów lub problemów związanych z samodzielnym wykonaniem bądź wysokim kosztem zakupu niezbędnych lamp, transformatorów głośnikowych czy zasilających. Opisany wzmacniacz, zbudowany z użyciem (jeszcze względnie łatwo dostępnych na aukcjach internetowych) tranzystorów germanowych AD161/162, może stać się pierwszym krokiem w technice retro, bez nadwyższania budżetu lub obawy o bezpieczeństwo.



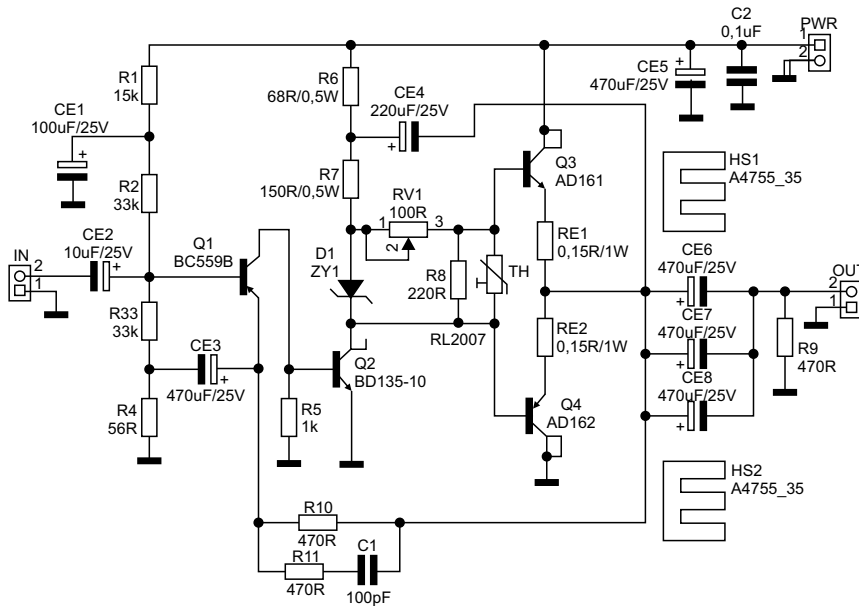
Zaprezentowany wzmacniacz mocy może być tranzystorowym odpowiednikiem konstrukcji opartej na lampach E(P) CL82 lub E(P)CL86 o mocy ok. 2...4 W. Schemat układu pokazano na **rysunku 1**.

Konstrukcja wzmacniacza bazuje na rozwiązaniu udostępnionym przez nieistniejącą już firmę Telefunken w piątym tomie „Informatora Radiowo-Warsztatowego”, w rozdziale „Wzmacniacze m.cz. z tranzystorami komplementarnymi” (rysunek 340, strona 236, WKŁ, wydanie z 1974 r.). Układ został wybrany ze względu na prostotę, łatwo dostępne elementy oraz brak transformatorów sterujących czy głośnikowych, przeznaczonych do układów tranzystorowych. Schemat został poddany niewielkim modyfikacjom, tak aby dostosować

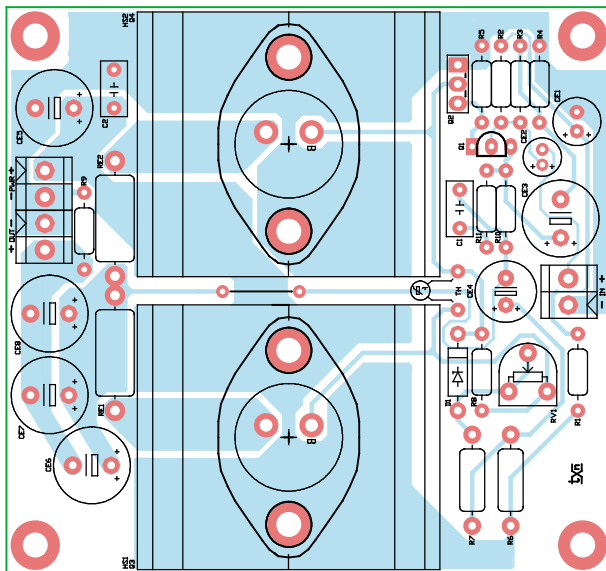
elementy do aktualnie dostępnych na rynku oraz dopasować poziomy sygnałów do standardów spotykanych we współczesnym sprzęcie audio. Wzmacniacz mocy składa się tylko z trzech stopni wzmocnienia. Sygnał wejściowy doprowadzony jest do złącza IN, gdzie – po separacji składowej stałej przez CE2 – trafia do stopnia wzmocnienia wstępnego z tranzystorem Q1. Blok ten jest sprzężony bezpośrednio z tranzystorem Q2 stopnia sterującego, odpowiadającego za wysterowanie komplementarnej końcówki mocy z tranzystorami Q3, Q4. CE4 jest elementem sprzężenia dynamicznego bootstrap. Tranzystory końcowe pracują z niewielkim prądem spoczynkowym ok. 10 mA, wpływającym na zmniejszenie zniekształceń skrośnych (klasa AB). Obwód

regulacji prądu spoczynkowego składa się z diody Zenera D1 o spadku napięcia ok. 0,7...0,8 V oraz dzielnika RV1, R8, TH i rezystorów emiterowych RE1,2.

Potencjometr RV1 odpowiada za ustalenie prądu spoczynkowego końcówki mocy, a termistor TH 50 Ω – za kompensację termiczną jego wartości. Jest to już rozwiązanie praktycznie niespotykane we współczesnych wzmacniaczach, gdzie kompensacja opiera się na zastosowaniu diod lub tranzystorów sprzężonych termicznie lub wręcz wykonanych w strukturze tranzystorów mocy. Należy pamiętać, że tranzystory germanowe mają znacznie niższe napięcie U_{be} (rzędu 0,2...0,3 V) w porównaniu do tranzystorów krzemowych (0,6...0,7 V), co wymusza nieco nietypowe



Rysunek 1. Schemat wzmacniacza



Rysunek 2. Rozmieszczenie elementów na płytce wzmacniacza

wartości elementów w obwodzie kompensacji. Jako termistor zastosowano model RL2007-32.8-59-D1 firmy Amphenol Thermometrics, dostępny u większości dużych dystrybutorów elementów elektronicznych. Nie należy lekceważyć roli tego niepozornego elementu, gdyż tranzystory germanowe mają znacznie węższy zakres dopuszczalnej temperatury pracy struktury – w przypadku AD161/162 wynosi on maksymalnie 90°C, a także silną zależność parametrów od temperatury, w porównaniu do tranzystorów krzemowych, co już przy niewielkiej mocy strat, wynoszącej tylko 4 W, może doprowadzić do szybkiego uszkodzenia w przypadku nieprawidłowej kompensacji. Tranzystory Q3, Q4 umieszczone są na radiatorach HS1, 2 wykonanych z typowego profilu A4755 o długości 35 mm. Montaż tranzystorów odbywa się bezpośrednio na radiatorze, bez podkładek izolacyjnych, a tylko na cienkiej warstwie pasty termoprzewodzącej w celu poprawienia kontaktu termicznego. Obudowy tranzystorów AD161/162 typu SOT-9 do złudzenia przypominają obudowy TO-66 polskich tranzystorów krzemowych BD254/255 lub BD354/355, ale różnią się kilkoma wymiarami i nie można bezpośrednio zastosować podkładek lub radiatorów do nich przeznaczonych. Do poprawnej kompensacji temperaturowej niezbędne jest umieszczenie termistora TH pomiędzy radiatorami oraz jego dodatkowe zaizolowanie kawałkiem cienkościennej rurki termokurczliwej wypełnionej pastą termoprzewodzącą. Należy pamiętać, że obudowy tranzystorów połączone są z kolektorami i w związku z tym na radiatorze Q3 występuje potencjał zasilania, a na Q4 – masa układu. Równolegle połączone kondensatory CE6, 7, 8 zapewniają separację składowej stałej

REKLAMA



INNOWACYJNE PRODUKTY
INNOWACYJNE TECHNOLOGIE

-  **Kontraktowy montaż elektroniki**
-  **Konwertowanie materiałów**
-  **Moduły laserowe**
-  **Szablony SMT**
-  **Dystrybucja**

**Wyznaczamy najwyższe
standardy jakości
w naszej branży**



Semicon Sp. z o.o. ul. Zwolenńska 43/43A, 04-761 Warszawa

22 615 73 71

info@semicon.com.pl

semicon.com.pl

Wykaz elementów:

Rezystory:

R1: 15 k Ω (1%, metalizowany)
 R2, R3: 33 k Ω (1%, metalizowany)
 R4: 56 k Ω (1%, metalizowany)
 R5: 1 k Ω (1%, metalizowany)
 R6: 68 Ω /0,5 W (1%)
 R7: 150 Ω /0,5 W (1%)
 R8: 220 Ω (1%, metalizowany)
 R9...R11: 470 Ω (1%, metalizowany)
 RE1, RE2: 0,15 Ω /1 W (1%)

RV1: potencjometr montażowy
 100 Ω (R=5 mm)

Kondensatory:

C1: 100 pF (50 V, foliowy, R=5mm)
 C2: 0,1 μ F (50 V, foliowy, R=5mm)
 CE1: 100 μ F/25 V (elektrolityczny f=6,3 mm, R=2,5 mm)
 CE2: 10 μ F/25 V (elektrolityczny f=5 mm, R=2 mm)

CE3, CE5...CE8: 470 μ F/25 V (elektrolityczny f=10 mm, R=5 mm)
 CE4: 220 μ F/25 V (elektrolityczny f=8 mm, R=3,5 mm)

Półprzewodniki:

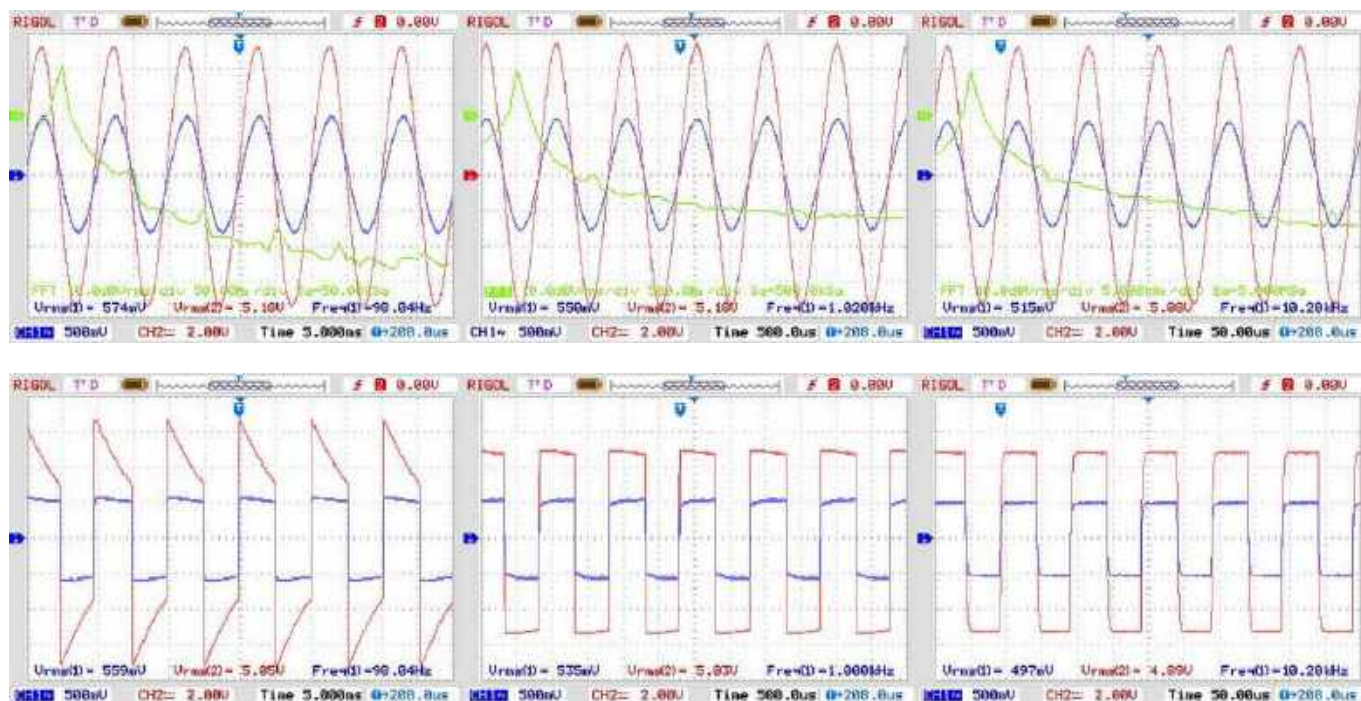
D1: dioda Zenera ZY1
 Q1: BC559B (TO-92)
 Q2: BD135-10 (TO-126)
 Q3: tranzystor germanowy AD161 (SOT-9,

parowany)

Q4: tranzystor germanowy AD162 (SOT-9, parowany)

Pozostałe:

IN, OUT, PWR: złącze DG 3,5 mm 2 pin (DG381-3.5-2)
 HS1, HS2: radiator aluminiowy (A4755_35)
 TH: termistor RL2007-32.8-59-D1 (Amphenol, 50 Ω)



Rysunek 3. Przykładowe przebiegi uzyskane podczas pomiarów wzmacniacza

od głośnika, ich wartość dobrana jest tak, aby zapewnić przenoszenie niskich częstotliwości od ok. 40 Hz (w przypadku głośnika 4 Ω). Dla poprawy parametrów wzmacniacz objęty jest sprzężeniem zwrotnym. Należy pamiętać, że układ nie jest wyposażony w zabezpieczenia przeciwzwarciowe i przeciążeniowe, zwarcie zacisków głośnika może zatem zakończyć się uszkodzeniem tranzystorów mocy. Głośnik podłączony jest do złącza OUT, a zasilanie – do gniazda PWR. Wzmacniacz pracuje poprawnie w zakresie 9...18 V przy obciążeniu 4 Ω , zalecam jednak nieco niższe napięcie 9...15 V, aby nie zwiększać strat w końcówce. Przy obciążeniu 8 Ω można korzystać z wyższego napięcia zasilania (15...18 V). Dobrze byłoby, aby napięcie pochodziło z zasilacza stabilizowanego, np. LM7815/LM7818. Osiągana moc, w zależności od warunków, zawiera się w zakresie 2...4 W przy zastosowaniu głośnika 4 Ω lub 2...3 W przy 8 Ω (dane dotyczą pasma przenoszenia 40 Hz...20 kHz i zniekształceń poniżej 5%).

Dla uproszczenia montażu i uruchomienia wszystkie elementy końcówki mocy,

włączając tranzystory mocy Q3, 4 wraz z ich radiatorami, umieszczone są na płytce drukowanej. Rozmieszczenie elementów pokazano na **rysunku 2**, a zmontowany moduł – na **fotografii tytułowej**.

Układ złożony ze sprawdzonych elementów działa po włączeniu zasilania, ale wymaga wyregulowania prądu spoczynkowego. Pierwsze uruchomienie należy wykonać z użyciem zasilacza laboratoryjnego z ograniczeniem prądowym ustawionym na 100 mA oraz przy obniżonym do 9 V napięciu zasilania. Do wyjścia wzmacniacza podłączamy rezystor 8 Ω /10 W, krótkim przewodem zwieramy wejście IN, potencjometr RV1 ustawiamy na minimum prądu spoczynkowego (zwarłe wyprowadzenia 1–2). Po włączeniu zasilania kontrolujemy w punkcie połączenia rezystorów emiterowych RE1, 2 obecność napięcia zbliżonego do połowy napięcia zasilania. Następnie przełączając woltomierz pomiędzy emiter tranzystorów mocy, ustawiamy prąd spoczynkowy na 10 mA, co odpowiada napięciu ok. 10,2 mV. Jeżeli wszystko działa

poprawnie, do wyjścia podłączamy oscyloskop, a do wejścia generator sinusoidalnego sygnału testowego. Zwiększając poziom sygnału testowego, na wyjściu powinniśmy obserwować czysty, niezniekształcony sygnał, a zwiększając napięcie zasilania do 15...18 V, możemy zmierzyć osiągalną moc wyjściową. Co kilka minut obserwujemy zachowanie się prądu spoczynkowego, który nie powinien zmienić się o więcej niż 20%. Sprawdzamy także temperaturę radiatorów tranzystorów – ta nie powinna przekraczać 50°C. Jeżeli tak nie jest, wykonujemy korekcję prądu potencjometrem RV1. Ostateczna regulacja prądu spoczynkowego musi zostać wykonana po ustaleniu się warunków termicznych płytki, czyli po ok. 30 minutach pracy ciągłej.

Przykładowe przebiegi otrzymane podczas pomiarów wzmacniacza zasilanego napięciem 18 V i obciążonego rezystorem 8 Ω zaprezentowano na **rysunku 3**.

Przyjemnego odsłuchu...

Adam Tatuś, EP

REKLAMA

www.facebook.com/ElektronikaPraktyczna



W ofercie AVT*

AVT6070

Najważniejsze parametry:

- wytwarzanie dźwięków o zmiennej częstotliwości,
- częstotliwość sygnału zmienia ruchem rąk wokół anteny oraz potencjometrem,
- wbudowany wzmacniacz współpracujący z głośnikiem o impedancji 8 Ω,
- zasilanie: 5 V (DC),
- pobór prądu: 10...200 mA.

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

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

Dodatkowe materiały do pobrania ze strony www.ulubionykiosk.pl/media

- | | |
|-------------|--|
| AVT5798 | Efekt do gitary z zastosowaniem cyfrowego przetwarzania sygnałów (EP 8/2020) |
| Projekt 239 | Multieffekt dla wokalistów VEP10 (EP 1/2019) |
| AVT5641 | Procesor wokalny z efektami echa DRP-10 (EP 10/2018) |
| AVT5576 | Moduł „delay/reverb” (EP 4/2017) |
| --- | Pogłos analogowy (EP 3/2017) |
| AVT5569 | Mikser Dry/Wet (EP 2/2017) |

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

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

Theremin – elektroniczny instrument

Jęki, skrzypienie, piski, a nawet wrzaski – to wszystko możemy wygenerować samodzielnie (i to analogowo) wyłącznie przy użyciu ruchów swoich rąk. W czasopiśmie dla elektroników po prostu nie mogło zabraknąć pierwszego, popularnego instrumentu elektronicznego! O czym mowa? Oczywiście o słynnym thereminie.

Minęło blisko 100 lat odkąd Lew Siergiejewicz Termen, znany również jako Leon Theremin, uzyskał amerykański patent na swoje urządzenie. Jego instrument muzyczny działał na zupełnie innej zasadzie niż te znane wcześniej, bowiem dźwięk był w nim generowany całkowicie elektronicznie. Przez dekady to urządzenie było rozwijane, a różni artyści znajdowali dla niego coraz to nowe zastosowania.

W tym artykule opisuję prostą, półprzewodnikową wersję theremina, którą za niewielkie pieniądze każdy może poskładać w swoim domu. Jednak jego zasada

działania pozostaje taka sama, jak oryginału z ubiegłego wieku, bowiem bazuje na odstrajaniu generatora przy zbliżaniu ręki do podłączonej doń... anteny.

Budowa

Schemat ideowy omawianego układu znajduje się na **rysunku 1**. Całość składa się z dwóch pracujących niezależnie generatorów zbudowanych na bramkach

NAND z wejściami Schmitta i wykonanych w technologii CMOS. Po zwarceniu obu wejść tworzą inwertery z histerezą, zatem do zbudowania pełnoprawnych generatorów wystarczy już tylko dwa elementy: rezystor tworzący pętlę sprzężenia zwrotnego, jak również kondensator, wprowadzający opóźnienie w przełączaniu wyjścia generatora.

Częstotliwość pracy pierwszego generatora, powstałego w oparciu na bramce US1A, może być zmieniana poprzez zbliżanie ręki do anteny przylutowanej



Fotografia 1. Zmontowany układ prototypowy

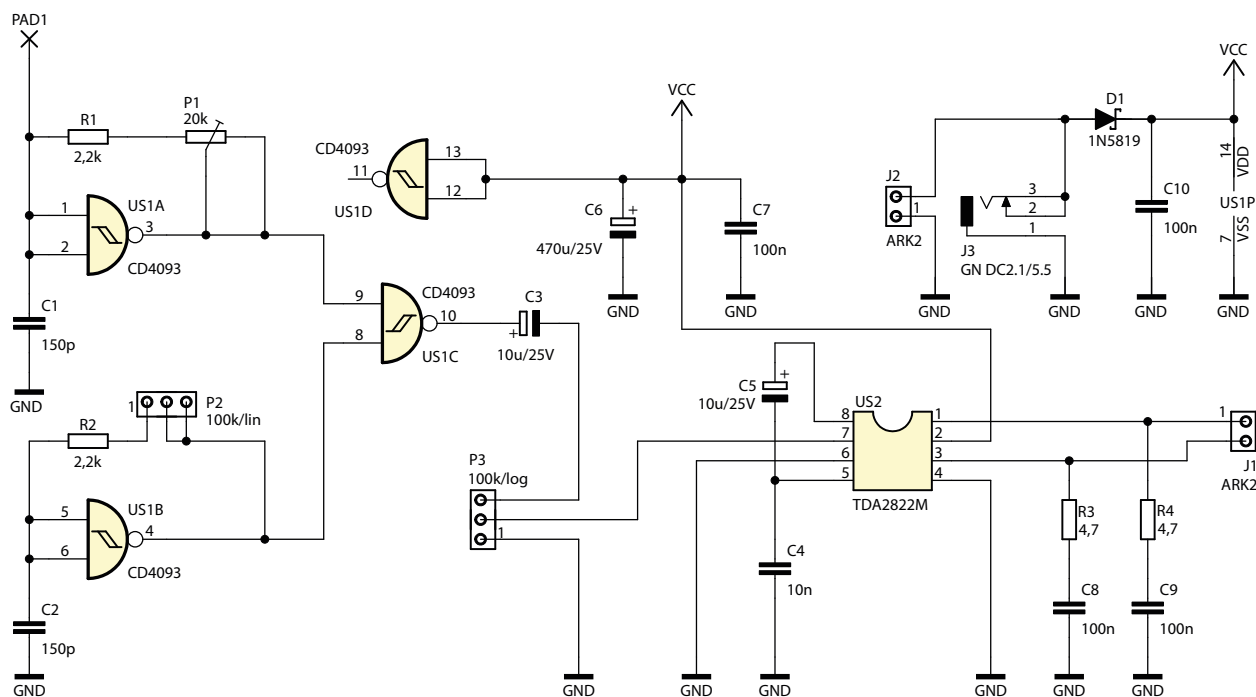
REKLAMA

LASEROWE SZABLONY DO MONTAŻU SMT

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



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



Rysunek 1. Schemat ideowy theremina

do pola lutowniczego PAD1. Wypadkowa rezystancja $R1+P1$ (gdzie $P1$ to regulowany opór potencjometru montażowego) ustala bazową częstotliwość pracy, natomiast pojemność kondensatora $C1$ jest zwiększana poprzez zbliżanie ręki do anteny – tworzy się wówczas dodatkowa pojemność między wejściami bramki $US1A$ a masą. Istotne jest, aby pojemność $C1$ była stosunkowo niska, ponieważ dodatkowa pojemność wynikająca ze zbliżania ręki to zaledwie kilka-kilkanaście pikofaradów. Zbyt wysoka pojemność $C1$ pozwalałaby na przestrajanie tego generatora jedynie w bardzo wąskim zakresie.

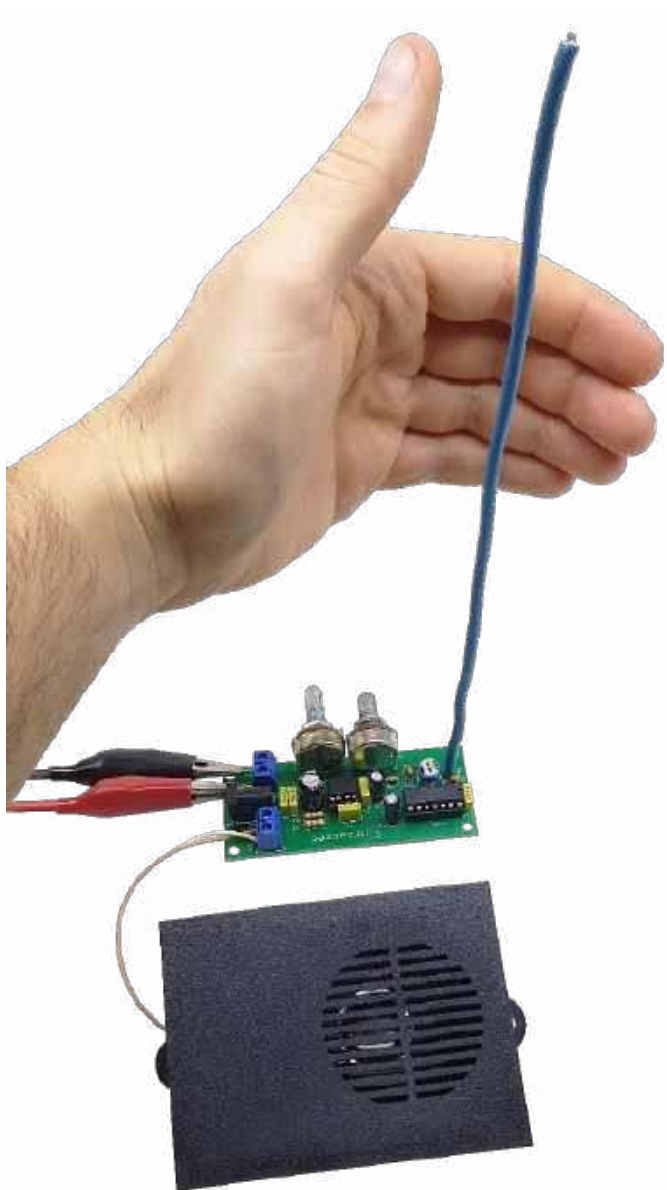
Drugi generator pracuje z ustaloną częstotliwością, którą można regulować potencjometrem (tym razem z gałką $P2$). Sygnały z obu tych generatorów są mnożone przez trzecią bramkę NAND, co można porównać do działania mieszacza w układzie heterodyny – tyle, że cyfrowego, dwustanowego. Dzięki temu sygnał wychodzący z bramki $US1C$ ma dwie składowe, zależne zarówno od generatora przestrajanego anteną, jak i od oscylatora o stałej częstotliwości (ustalonej potencjometrem $P2$), który w tym układzie możemy przyrównać do heterodyny.

Składowa stała otrzymanego sygnału jest oddzielana przy użyciu kondensatora $C3$. Potencjometr $P3$ służy do regulacji amplitudy składowej zmiennej, która z kolei trafia na wzmacniacz typu TDA2822, pracujący w układzie mostkowym. Dzięki takiej topologii jest on w stanie wysterować niewielki głośnik. Składowe o wysokiej częstotliwości, których na wejście układu $US2$ trafia bardzo dużo, są tłumione przez rezystory $R3$ i $R4$, pracujące w szeregowych obwodach RC (tak zwanych obwodach Zobla).

Montaż i uruchomienie

Układ został zmontowany na jednostronnej płytce drukowanej o wymiarach $80\text{ mm} \times 35\text{ mm}$. Jej wzór ścieżek oraz schemat montażowy pokazano na **rysunku 2**. W odległości 3 mm od krawędzi płytki znalazły się cztery otwory montażowe, każdy o średnicy 3,2 mm.

Montaż proponuję przeprowadzić w sposób typowy, czyli rozpoczynając od elementów o najmniejszej wysokości obudowy, w tym wypadku rezystorów i diod półprzewodnikowych. Pod układy scalone proponuję zastosować podstawki, aby ułatwić ich wymianę w razie uszkodzenia. Na samym końcu polecam wlutować antenę, którą można wykonać z odcinka izolowanego drutu o przekroju $1,5\text{ mm}^2$ i długości ok. 25 cm – przy takim przekroju będzie odpowiednio sztywna. Zmontowany układ można zobaczyć na **fotografii 1**.



Fotografia 2. Przykład użycia

Wykaz elementów:

Rezystory:

R1, R2: 2,2 k Ω (THT, 0,25 W)
 R3, R4: 4,7 Ω (THT, 0,25 W)
 P1: 20 k Ω (potencjometr montażowy leżący jednoobrotowy)
 P2: 100 k Ω (potencjometr liniowy na panel)
 P3: 100 k Ω (potencjometr logarytmiczny na panel)

Kondensatory:

C1, C2: 150 pF (THT, R=5 mm, monolityczny – opis w tekście)

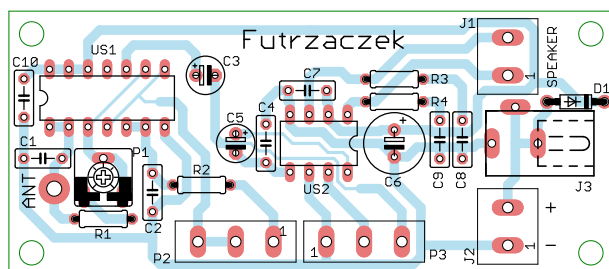
C3, C5: 10 μ F/25 V (THT, R=2,5 mm, elektrolityczny)
 C4: 10 nF (THT, R=5 mm, MKT)
 C6: 470 μ F/25 V (THT, R=3,5 mm, elektrolityczny)
 C7...C10: 100 nF (R=5 mm, MKT)

Półprzewodniki:

D1: 1N5819
 US1: CD4093 (DIP14)
 US2: TDA2822M (DIP8)

Pozostałe:

J1, J2: złącze śrubowe ARK2/500
 J3: gniazdo DC 2,1/5,5 mm
 Podstawka DIP8
 Podstawka DIP14
 Odcinek izolowanego drutu o przekroju 1,5 mm² i długości ok. 25 cm (opis w tekście)



Rysunek 2. Schemat montażowy i wzór ścieżek płytki

Prawidłowo złożony układ nie wymaga uruchamiania i jest od razu gotowy do pracy. Do złącza J1 należy podłączyć głośnik o impedancji 8 Ω lub większej i mocy co najmniej 1 W. Z kolei zasilanie układu można podłączyć do złącza J2 (odizolowane przewody) lub J2 (gotowy zasilacz wtyczkowy z wyjściem DC 2,1/5,5 mm). Pomimo że oba użyte w projekcie układy scalone mogą być zasilane napięciem z szerokiego przedziału, według moich doświadczeń najlepiej nadaje się wartość 5 V. Jest to kompromis pomiędzy mocą wydzielaną w układzie US a maksymalną głośnością generowanych przez układ dźwięków (przy jednoczesnym zachowaniu wysokiej

wrażliwości na odstranianie). Pobór prądu przy takim napięciu zasilającym waha się od 10 mA do 200 mA, zależnie od aktualnego poziomu dźwięku.

Układ w akcji jest widoczny na **fotografii 2**. Po podłączeniu głośnika i zasilania, z tego pierwszego powinien wydobywać się z grubsza jednostajny ton. Zbliżając i oddalając ręce od anteny, można zmieniać częstotliwość i charakter sygnału. Potencjometrem P1 ustala się bazową częstotliwość oscylacji generatora przestrajanego ręką, zaś za pomocą P2 – częstotliwość generatora „heterodyny”. Układ jest bardzo podatny na modyfikacje – można wymienić zarówno kondensatory C1 i C2 na elementy o innej pojemności (im większa pojemność, tym niższa częstotliwość), jak też zmienić rozmiary lub kształt anteny. Można też zmienić potencjometry P1 i P2 na elementy o wyższej rezystancji. W bazowej wersji tego układu oscylator z kondensatorem C1 uzyskuje częstotliwości z przedziału 380 kHz... 1,6 MHz, zaś z kondensatorem C2 od 105 kHz do 1,6 MHz. Słyszalne są zatem jedynie te składowe, które stanowią wynik zmieszania się tych częstotliwości i ich harmonicznych.

Życzę Czytelnikom udanego komponowania własnych dźwięków!

Michał Kurzela, EP

REKLAMA

UWAGA! Tylko prenumeratorzy czasopism „Elektronika dla Wszystkich”, „Elektronika Praktyczna”, „Świat Radio” oraz „Elektronik” mogą korzystać z atrakcyjnych rabatów w Sklepie AVT:

- ✓ do 50% na wydania specjalne czasopism Wydawnictwa AVT
- ✓ 20% na kity w wersji A (płytki drukowane do projektów AVT)
- ✓ 10% na pozostałe wersje kitów: (A+, B, C, D)
- ✓ 10% na książki
- ✓ 5% na pozostałe produkty z oferty sklepu

Ponadto każdy prenumerator ww. czasopism korzysta z rabatów od 30% do 50% na zakup czasopism z oferty www.UlubionyKiosk.pl

K L U B
AVT
ELEKTRONIKA

Jak uzyskać rabat? Podczas zamówienia powołaj się na swój numer prenumeraty – otrzymasz go mailowo po zakupie prenumeraty wraz z kartą członkowską Klubu AVT-Elektronika.

**Najważniejsze parametry:**

- pasmo: KF 7 MHz (7085...7185 kHz)/40 m,
- pośrednia przemiana częstotliwości,
- odbiór i nadawanie sygnałów modulowanych jednostęgowo (SSB),
- współpraca z zestawem słuchawkowo-mikrofonowym wyposażonym w przycisk PTT,
- zasilanie: pojedyncze 12 V (DC)/500 mA,
- PCB pasująca do obudowy Z78.

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

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

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

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

W ofercie AVT*

AVT6066

TRX Ewa 40 m – transceiver QRP początkującego krótkofalowca (1)

Pomimo łatwej dostępności fabrycznych transceiverów KF o mocy rzędu 100 W, transceivery małej mocy (QRP) są chętnie konstruowane oraz zabierane na wakacje czy urlopy. Praca z niewielką mocą, na własnoręcznie wykonanym, prostym i małogabarytowym urządzeniu zasilanym z akumulatora daje dużo przyjemności, nie tylko początkującym krótkofalowcom.

W ofertach handlowych, głównie w sklepach internetowych, można znaleźć wiele układów nadawczo-odbiorczych do samodzielnego montażu, bardzo często skomplikowanych i opartych na drogich elementach, w tym różnych, nowoczesnych układach scalonych. Dostępne są też gotowe moduły, z których można złożyć transceiver, ale nie mają one większych walorów dydaktycznych.

Prezentowany transceiver, na jedno z podstawowych pasm częstotliwości amatorskich (7 MHz), potocznie nazywanych czterdziestką, powstał z myślą o początkujących, licencjonowanych radioamatorach, którzy chcą zbudować bardzo proste oraz tanie urządzenie nadawczo-odbiorcze o niewielkich wymiarach i poznać przy tym podstawy radiotechniki.

Przy projektowaniu układu autorowi przyświecały dwa znane przysłowia:

- próżnym działaniem jest zbudowanie czegoś z dużej liczby elementów, jeśli można je wykonać z mniejszej ich liczby;
- największy sukces zostaje osiągnięty, gdy nie można już zmniejszyć liczby elementów.

Układ został ograniczony do niezbędnego minimum z użyciem popularnych



podzespołów i minimalnej liczby nawijanych samodzielnie cewek, przy zachowaniu dobrych parametrów urządzenia.

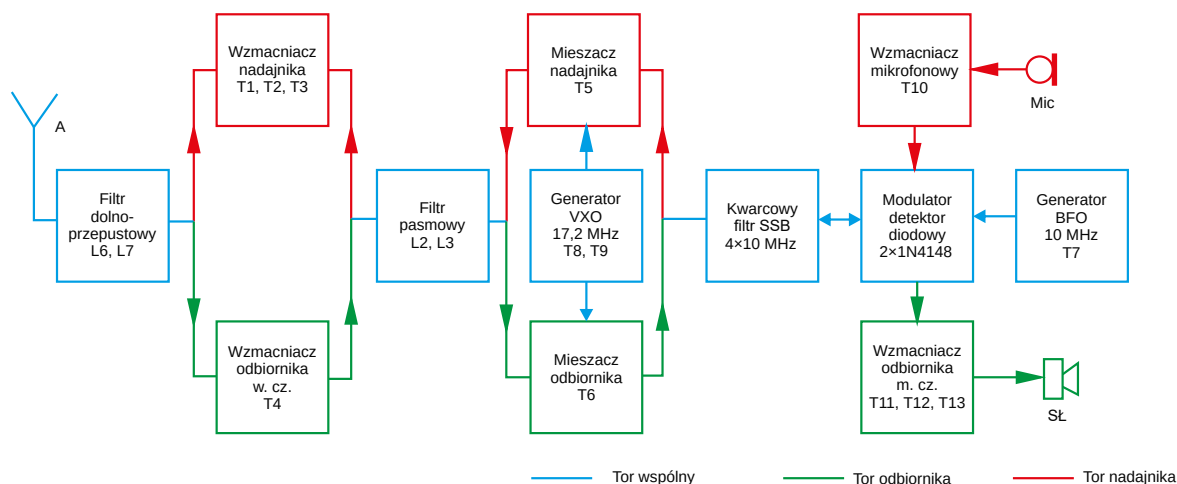
TRX pracuje w układzie z pośrednią przemianą częstotliwości i umożliwia odbiór oraz nadawanie sygnałów fonicznych – jednostęgowych (LSB), w najczęściej zajmowanym przez polskie stacje wycinku pasma 40 m. Umożliwia współpracę z anteną na pasmo 40 m, zasilaną kablem koncentrycznym, a także z zestawem słuchawkowo-mikrofonowym wyposażonym w przycisk PTT.

Prezentowany układ różni się zasadniczo od wszystkich aktualnie dostępnych rozwiązań, ponieważ został wykonany z użyciem jedynie 13 popularnych tranzystorów i jednego stabilizatora, czyli łatwo dostępnych elementów, jakie autor miał pod ręką.

Moc wyjściowa nadajnika przy zasilaniu z zasilacza 13,8 V wynosiła około 3...4 W, zaś wymiary urządzenia to 85 × 154 × 43 mm.

Podobne układy na dyskretnych elementach były budowane już dawno, kiedy na rynku jeszcze nie było układów scalonych, ale opisywane rozwiązanie prezentuje nowe podejście do konstrukcji w oparciu o aktualnie dostępne podzespoły oraz wieloletnie doświadczenia autora w konstruowaniu bardziej skomplikowanych układów (Bartek, Antek, Junior, Zuch, Jędreka...).

Przystępując do budowy tego TRX-a, autor miał świadomość, że trudno jest skonstruować radio o niewielkich wymiarach, do którego części będą kosztować kilkadziesiąt złotych, a które będzie jednocześnie miało parametry porównywalne z transceiverami fabrycznymi. Wprawdzie można zrobić SDR-a (jednak zazwyczaj



Rysunek 1. Schemat blokowy transceivera

potrzeba jeszcze komputera) lub zaangażować do pracy jakiś nowoczesny syntezer, DDS i PIC, ale wtedy to już nie będzie proste i tanie radio home made, tylko powstanie konstrukcja droższa, odpowiednia dla bardziej zaawansowanych.

Choć w urządzeniu są liczne uproszczenia układowe i ograniczona liczba zastosowanych elementów, to jednak uzyskane parametry umożliwiają w sprzyjających warunkach propagacyjnych normalną pracę foniczną w wycinku pasma 40 m.

Transceiver ten można polecić początkującym krótkofalowcom nie tylko ze względu na niską cenę użytych podzespołów, ale przede wszystkim z uwagi na walory edukacyjne. W takim „szkolnym” urządzeniu bardzo łatwo jest prześledzić tory sygnałów oraz poznać pracę poszczególnych bloków i bez problemu zrozumieć zasadę działania transceivera.

Konstrukcja ta może być pierwszą wprawką, po zdobyciu licencji, do budowy bardziej skomplikowanych układów nadawczo-odbiorczych i poznawania klasycznej radiotechniki oraz tajników krótkofalarstwa na popularnym paśmie HF.

Schemat blokowy transceivera pokazano na **rysunku 1**.

Jak łatwo się zorientować, urządzenie pracuje w układzie z pojedynczą przemianą częstotliwości, z podwójnym użyciem filtrów (dolnoprzepustowego, pasmowego i kwarcowego SSB), generatorów VXO i BFO oraz diodowego modulatora-detektora.

Dwukierunkowe zastosowanie filtrów oraz modulatora-detektora podczas nadawania i odbioru znacznie uprościło konstrukcję urządzenia (wyeliminowało konieczność komutacji sygnału).

Wybór częstotliwości pośredniej 10 MHz wynika głównie z użycia w układzie VXO rezonatora kwarcowego 17,2 MHz (zamiast tradycyjnego VFO z przestrajającym obwodem LC).

Kompletny schemat ideowy urządzenia prezentuje **rysunek 2**. Konstrukcja została

uproszczona do minimum, w tym z użyciem głównie jednego typu popularnych tranzystorów typu BC547B lub podobnych NPN. Obecnie produkowane tranzystory NPN małej mocy są bardzo tanie i choć katalogowo przeznaczone dla m.c.z., oferują bardzo wysoką częstotliwość graniczną, ponad 150 MHz, co umożliwia zastosowanie ich w układach radiowych KF.

Mając na uwadze mniej doświadczonych elektroników, najpierw zostanie omówiona zasada działania transceivera podczas odbioru (RX), a potem podczas nadawania (TX).

Odbiornik – RX

Sygnał antenowy po przejściu poprzez filtr dolnoprzepustowy L7, L6, C1...C3 jest podany – przez styki przekaźnika – na tłumik wejściowy w postaci potencjometru POT1. W ten prosty sposób, poza regulacją siły głosu, został rozwiązany problem z obniżaniem sygnału wejściowego (przesterowaniem układu silnymi sygnałami pochodzącymi np. od pobliskiej stacji sąsiada krótkofalowca).

Ograniczony do 8 MHz sygnał jest skierowany na wzmacniacz z tranzystorem T4. Pracuje on w układzie OB, dzięki czemu ma niską impedancję wejściową, porównywalną z filtrem dolnoprzepustowym oraz dużą impedancję wyjściową, nieobciążającą obwodu rezonansowego L2–C16. Wzmacniacz pracuje tylko podczas odbioru i jest załączany napięciem 12 V/+RX poprzez diodę D2.

Wzmocniony sygnał z anteny, poprzez kondensator sprzęgający C18, dochodzi do drugiego obwodu rezonansowego L3–C19, który pracuje razem z poprzednim obwodem jako dwupasmowy filtr środkowoprzepustowy, zarówno podczas odbioru, jak i nadawania.

Odfiltrowany sygnał w paśmie 40 m, poprzez kondensator C22, jest następnie kierowany na mieszacz odbiornika, zbudowany

w oparciu na tranzystorze T6. Do emitera tego tranzystora jest doprowadzany sygnał z generatora VXO. Taki obwód to najprostsz mieszacz, który jednocześnie wzmacnia sygnał. Pewnym mankamentem tego układu są wyjściowe, niepożądane produkty mieszania, ale zostały one w dostatecznym stopniu odfiltrowane w dalszej części toru sygnałowego. Układ pracuje tylko podczas odbioru i jest załączany napięciem 12 V/+RX poprzez diodę D4 (w podobny sposób jest rozwiązany mieszacz nadajnika).

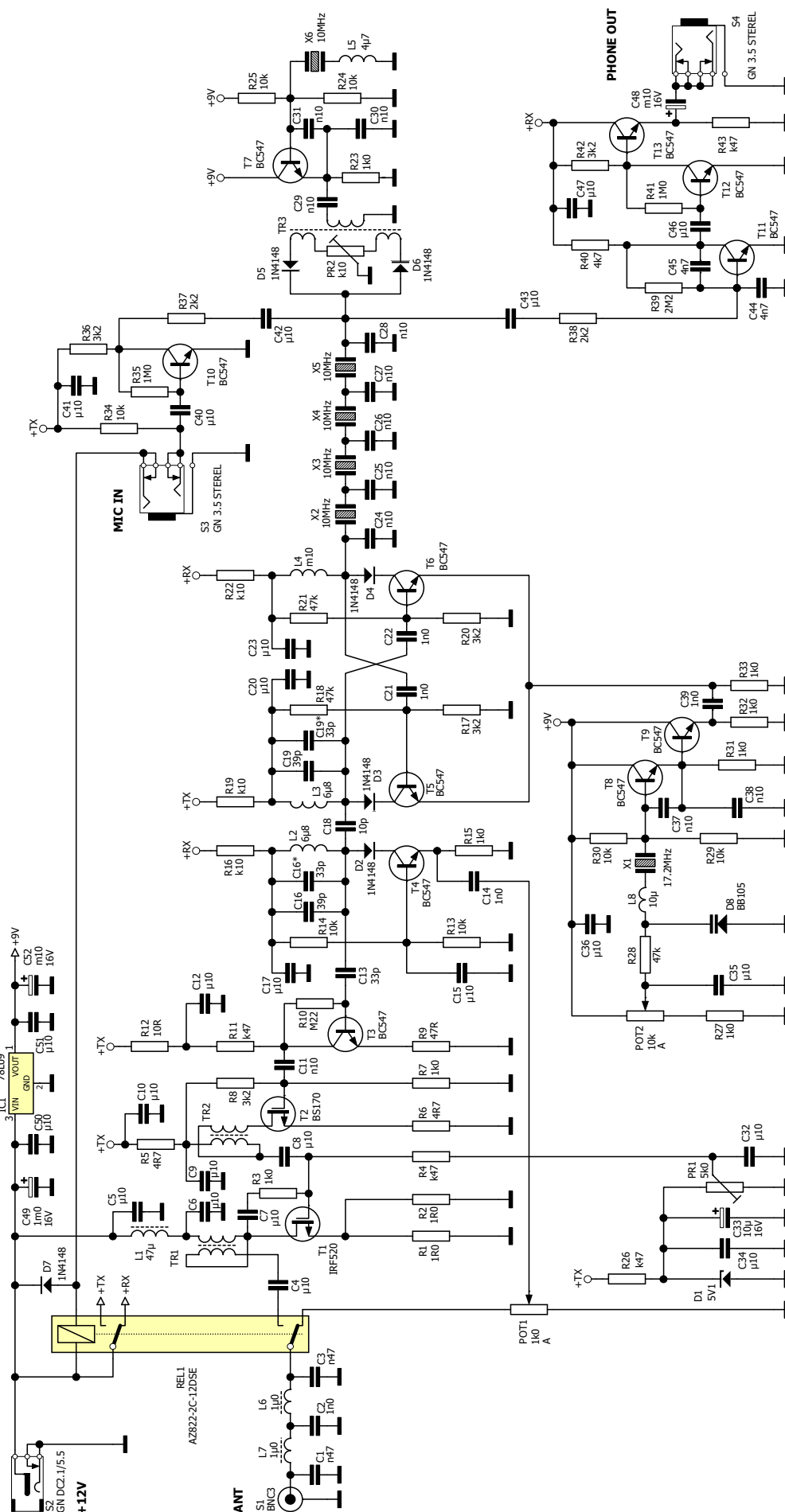
Wydzielony w obwodzie kolektora sygnał pośredniej częstotliwości (różnica częstotliwości sygnałów z generatora i wejściowego) jest skierowany na filtr kwarcowy SSB, zestawiony w układzie drabinkowym z czterech rezonatorów o częstotliwości 10 MHz. Przy zastosowaniu czterech jednakowych rezonatorów kwarcowych i współpracujących kondensatorów C24...C28 (po 100 pF każdy) filtr zapewnia pasmo przenoszenia około 2,5...2,7 kHz. Bezpośrednio po filtrze (z pominięciem wzmacniacza p.c.z.) jest włączony zrównoważony modulator dwudiodowy D5–D6.

Pojawienie się sygnału p.c.z. na wyjściu filtru kwarcowego powoduje „równoważenie” układu i dalsze przejście sygnału z generatora kwarcowego BFO.

Częstotliwość tego generatora leży na dolnej części charakterystyki filtru kwarcowego p.c.z., co jest niezbędne do odtworzenia brakującej wstęgi bocznej sygnału wejściowego.

Częstotliwość pracy generatora z tranzystorem T7 zależy głównie od rezonatora kwarcowego oraz włączonej w szereg z nim cewki (dławik L5 o indukcyjności 4,7 μH), zapewniającej obniżenie częstotliwości BFO o ponad 1 kHz. Między BFO a detektorem/modulatorem może być włączony stopień separujący, ale w praktyce nie było takiej potrzeby (w przeciwieństwie do VXO, gdzie taki zabieg okazał się przydatny).

W wyniku zmieszania sygnału p.c.z. z sygnałem wewnętrznego generatora



Rysunek 2. Schemat ideowy układu

uzyskuje się na wyjściu czytelny sygnał małej częstotliwości (0,3...3 kHz), który jest podawany poprzez dwójnik R38–C43 na wzmacniacz małej częstotliwości. W torze tym pracuje podwójny stopień wzmacniacza w układzie OE z tranzystorami T11 i T12. Trzeci tranzystor – T13 – nie daje wzmocnienia napięciowego, lecz silne wzmocnienie prądowe, ponieważ pracuje w układzie OC i służy do dopasowania niskiej impedancji słuchawek.

Kształtowanie charakterystyki sygnału m.c.z. w zakresie 0,3 kHz zapewniają kondensatory sprzęgające, zaś za ograniczenie powyżej 3 kHz odpowiada kondensator C45 w pętli ujemnego sprzężenia zwrotnego. Dodatkowy kondensator filtrujący C44 również ma wpływ na ograniczenie sygnału od strony wyższych częstotliwości.

Urządzenie jest przewidziane do współpracy z zestawem multimedialnym (mikrofon elektretowy + słuchawki niskoomowe).

Częstotliwość pracy generatora V XO z tranzystorem T8 jest ustalana elektronicznie poprzez napięcie zasilania podawane z potencjometru POT2 na diodę pojemnościową D8 typu BB105. Przy maksymalnym napięciu zasilania (9 V suwak w górnym położeniu) uzyskano na wyjściu częstotliwość 17,185 MHz, zaś przy niskim – częstotliwość około 7,085 MHz, czyli w efekcie uzyskano szerokość pasma V XO 100 kHz, co odpowiada najbardziej interesującemu wyinkowi pasma SSB od około 7000 do 7185 kHz.

Cały sekret tak szerokiego zakresu przestrajania generatora kwarcowego leży w doborze szeregowo włączonej z rezonatorem kwarcowym indukcyjności L8 oraz pojemności warikapu D8 (stosunkowo duża indukcyjność dławika oraz mała pojemność diody). Pomimo uproszczeń, układ generatora pracuje wyjątkowo stabilnie i co najważniejsze, od razu w docelowym zakresie 40 m, z reguły nie

Wykaz elementów:**Półprzewodniki:**

D1: 5V1 (4V7, 6V8)
D2...D7: 1N4148
D8: BB105
T1: IRF520
T2: BS170
T3...T13: BC547

Kondensatory:

C1, C3: 470 pF
C2, C14, C21, C22, C39: 1 nF
C4...C10, C12, C15, C17, C20, C23, C32, C34...C36, C40...C43,
C46, C47, C50, C51: 100 nF
C13: 33 pF
C16, C19: 72 pF (39 pF+33 pF)
C11, C24, C25...C31, C37, C38: 100 pF
C18: 10 pF
C33: 10 µF/16 V
C44, C45: 4,7 nF
C48, C52: 100 µF/16 V
C49: 1000 µF/16 V

Rezystory:

R1, R2: 1 Ω
R3, R7, R15, R23, R27, R31...R33: 1 kΩ
R4, R11, R26, R43: 470 Ω
R5, R6: 4,7 Ω
R8, R17, R20, R36, R42: 3,2 kΩ
R9: 47 Ω
R10: 220 kΩ
R12: 10 Ω
R13, R14, R24, R25, R29, R30, R34: 10 kΩ
R16, R19, R22: 100 Ω
R18, R21, R28: 47 kΩ
R37, R38: 2,2 kΩ
R39: 2,2 MΩ
R40: 4,7 kΩ
R35, R41: 1 MΩ
POT1: 1 kΩ/A (potencjometr obrotowy)
POT2: 10 kΩ/A (potencjometr obrotowy)
PR1: 5 kΩ (potencjometr montażowy)
PR2: 100 Ω (potencjometr montażowy)

Pozostałe:

REL1: przełącznik 12V (AZ822-2C-12DSE lub podobny, np. HFD27/012-S, V23042-A 2003-B20)
ANT (S1): gniazdo antenowe BNC do druku
MIC IN (S3): GN 3,5 STEREL (gniazdo słuchawkowe stereo mini jack)
PHONE OUT (S4): GN 3,5 STEREL (gniazdo słuchawkowe stereo mini jack)
+12V (S2): GN DC2.1/5.5 (gniazdo zasilania 12V)
Radiator: DY-AM (32×20×30 mm) lub podobny
Obudowa plastikowa Z78 (85×154×43 mm)
Pokrętła na osie potencjometrów 6 mm (2 szt.)
L1: 47 µH (10 zwojów DNE 0,6 na rdzeniu FT37-43)
L2, L3: 6,8 µH (dławiki osiowe)
L4: 100 µH (dławik osiowy)
L5: 4,7 µH (dławik osiowy)
L6, L7: 1 µH (16 zwojów DNE 0,4 na rdzeniu T37-2)
L8: 10 µH (dławik osiowy)
TR1, TR2: 2×10 zwojów DNE 0,4 na rdzeniu FT37-43
TR3: 3×10 zwojów DNE 0,2 na rdzeniu FT37-43
X1: 17,2 MHz
X2...X5: 10,0 MHz

wymagając korekty. Tranzystor T9 jest włączony w układzie OC i pełni funkcję separatora generatora.

Do zasilania generatorów zastosowano napięcie 9 V pochodzące ze stabilizatora UC1 typu 78L09.

Nadajnik – TX

Po naciśnięciu przycisku PTT przełącznik RE1 przełącza urządzenie z odbioru na nadawanie. Jedna para styków służy do przełączania anteny, a druga przełącza napięcie zasilania (odłącza szynę +12 V od odbiornika i doprowadza ją do nadajnika).

Sygnał z mikrofonu elektretowego, po wzmocnieniu w układzie z tranzystorem T10, jest podawany przez dwójnik R37–C42 na wejście modulatora diodowego. Pojawienie się sygnału m.cz. powoduje „rozwórnoważenie” układu mieszacza i podanie sygnału z generatora kwarcowego na wejście filtra kwarcowego.

Ponieważ częstotliwość generatora fali nośnej (BFO) leży na dolnej części charakterystyki filtra kwarcowego p.cz., na wyjściu filtra uzyskuje się sygnał z górną wstęgą boczną, zaś wstęga dolna (USB) jest odfiltrowana.

Sygnał USB z wyjścia filtra jest podany poprzez kondensator C21 na mieszacz nadajnika z tranzystorem T5. Mieszacz ten pracuje tylko podczas nadawania i jest załączany napięciem 12 V/+TX poprzez diodę D3. Podobnie jak w mieszaczu odbiornika, na emiter tranzystora T5 jest skierowany sygnał z generatora VXO (przestrajanego w wyżej podanym zakresie). Wyjściowy filtr dwuobwodowy L3–C19, L2–C16 filtruje i przepuszcza sygnał nadajnika w zakresie 7085...7185 kHz. Ponieważ częstotliwość wyjściowa emitowanego sygnału jest odejmowana od wyższej częstotliwości VXO, w efekcie uzyskuje się odwróconą wstęgę boczną sygnału, czyli dolną (LSB), jaka jest wymagana w paśmie 40 m (do 10 MHz LSB).

W układzie wzmacniacza nadajnika pracują 3 tranzystory. W stopniu mocy

i driverze zostały zastosowane tranzystory MOSFET z kanałem N (IRF520 i BS170).

Wprawdzie IRF520 jest przewidziany do pracy impulsowej, ale równie dobrze spisuje się także we wzmacniaczach mocy na niższych zakresach HF.

Pierwszy stopień nadajnika to klasyczny wzmacniacz OE z tranzystorem T3, a po nim jest driver T2 z tranzystorem unipolarnym BS170.

Sygnał SSB – bezpośrednio z drivera – jest kierowany przez obniżający transformator dopasowujący TR2 (przekładnia 4:1) do stopnia końcowego mocy z tranzystorem T1 (IRF520).

Dopasowanie obwodów drenu tranzystora do filtra dolnoprzepustowego zrealizowano za pomocą podwyższającego transformatora TR1 (1:4).

Na wyjściu znajduje się podwójny filtr dolnoprzepustowy eliminujący sygnały pozapasmowe (harmoniczne).

Jak wiadomo, jednym z najważniejszych parametrów wzmacniacza nadajnika SSB jest jego liniowość. Poprawną pracę drivera z tranzystorem T2 zapewnia dzielnik rezystorowy R7–R8 (prąd spoczynkowy około 40 mA), a w przypadku stopnia końcowego mocy z tranzystorem

T1 osiągnięto ją poprzez ustawienie punktu pracy za pomocą potencjometru PR1 (prąd spoczynkowy około 120 mA).

Aby uniezależnić punkt pracy tranzystora T1 od wartości napięcia zasilania, w układzie polaryzacji bramki stopnia końcowego jest włączona dioda Zenera D1, zasilana napięciem +TX.

Rezystory włączone w źródłach tranzystorów polowych wprowadzają niewielkie, ujemne sprzężenie zwrotne, wpływające pozytywnie na liniowość układu i kompensację temperaturową stopnia.

Urządzenie najlepiej jest zasilac z zewnętrznego źródła DC 12 V (akumulatora 12 V lub dobrze filtrowanego i stabilizowanego zasilacza transformatorowego), ponieważ w praktyce większość tanich zasilaczy impulsowych wprowadza spore zakłócenia. Maksymalny pobór prądu przy zasilaniu napięciem 13,8 V podczas nadawania dochodził do 500 mA, a podczas odbioru wynosił około 40 mA.

Układ modelowy był zasilany z powodzeniem także z trzech ogniw akumulatorów Li-Ion 3,7 V połączonych w szereg, ale moc nadajnika spadła do około 2 W.

Andrzej Janeczek SP5AHT
sp5aht@swiatradio.pl

REKLAMA

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

WWW.AKSOTRONIK.COM.PL

Aksotronik
ELEMENTY ELEKTRONICZNE

Magnesy neodymowe oraz ferrytowe
Ceny od 6,10zł

Przełączniki klawiszowe wodoodporne/płaskie
Ceny od 2,40zł

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

Przewodniki do przewodów
Ceny od 11,00zł

Kaski elektryczne zardzewiałe
Ceny od 0,33zł

Szeroki węgiel do elektronarzędzi
Ceny od 2,60zł

Przełączniki do elektronarzędzi zwykłe i elektronarzędziowo
Ceny od 7,00zł

Podłoża/organizery
Ceny od 0,95zł

Złącza termiczne Superseal
Ceny od 1,10zł

Zestawy drabek M2, M3 z nakrętkami i podkładkami
Ceny od 2,50zł

Uwaga!!! Powyższe ceny dotyczą zakupów minimalnych ilości hurtowych, poprzez nasz sklep internetowy.

W swojej ofercie posiadamy m.in.: półprzewodniki, diody, układy scalone, tranzystory, triaki, elementy optoelektroniczne, elementy dysznowe, łączniki, przełączniki, elementy akustyczne, rezystory, kondensatory, kwarce, podstawniki, moduły Arduino.

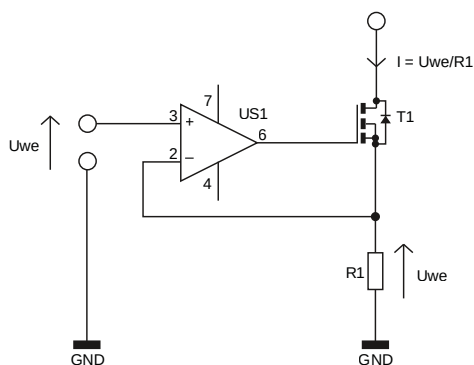
Zapraszamy do kontaktu: **INFO@aksotronik.com.pl**, tel: (22) 783-20-51

Źródło prądowe (prawie) idealne

Źródła prądowego wprawdzie nie kupimy w sklepie ze źródłami prądowymi, ale możemy je sobie zbudować. Na ten temat powstało wiele opracowań, ale jedno z często stosowanych rozwiązań układowych potrafi przysporzyć problemów. O czym mowa i jak owym trudnościom zaradzić?

Źródła prądowe da się zrealizować na wiele sposobów, różniących się przede wszystkim dokładnością ustalania natężenia prądu oraz możliwością jego regulacji. Nie będę się rozpisywał o wszystkich potencjalnych wariantach (począwszy od tranzystora bipolarnego pracującego jako mnożnik prądu bazy), gdyż takich układów w praktyce się nie spotyka. Natomiast bardzo często mamy styczność ze źródłem prądowym, które w literaturze określa się mianem precyzyjnego – i które może za takie uchodzić, jeżeli akurat zadziała.

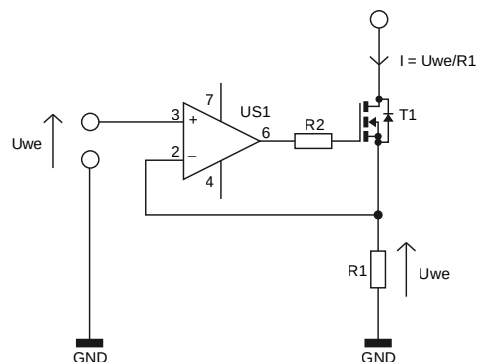
Mowa o układzie z **rysunku 1** bądź jakiegś jego zmodyfikowanej odmianie. W teorii układ ten bardzo dobrze zachowuje się w roli źródła prądowego. Wzmacniacz operacyjny US1 tak steruje bramką tranzystora T1, aby spadek napięcia na rezystorze R1 był równy napięciu wejściowemu U_{we} . Układ „wciąga” drenem T1 prąd o natężeniu równym $U_{we}/R1$, czego na bieżąco pilnuje wzmacniacz operacyjny. Taki sam układ, tylko z tranzystorem MOSFET z kanałem typu P, byłibyśmy w stanie zbudować dla prądu wypływającego od strony dodatniej szyny zasilania.



Rysunek 1. Podstawowy układ precyzyjnego źródła prądowego z tranzystorem MOSFET

Dlaczego w roli elementu wykonawczego znalazł się tranzystor unipolarny? Odpowiedź nie jest trudna: chodzi o praktycznie zerowy prąd bramki, gwarantujący równość prądu źródła (przez które przepływa prąd wywołujący spadek na R1) oraz prądu drenu (którego faktycznie używamy w innych częściach układu). Jedynie, co stoi na przeszkodzie w uzyskaniu faktycznie idealnego wyniku, to offset napięciowy wzmacniacza operacyjnego i tolerancja wykonania rezystora R1. W dalszych rozważaniach można się pokusić nawet o przeanalizowanie znaczenia prądów wejściowych wzmacniacza operacyjnego, zwłaszcza przy małym natężeniu wymuszanego prądu I. Tyle że to wszystko nie zda się na nic, jeżeli układ po prostu nie zadziała.

Dlaczego? Ponieważ najnormalniej w świecie się wzbudzi (choć nie zawsze, raczej w 99% przypadków). Dzieje się tak, ponieważ wzmacniacz operacyjny steruje niemal czystą pojemnością o wartości od kilku do kilkudziesięciu nanofaradów. Taki układ, objęty sprzężeniem zwrotnym, ma tendencję do popadania w oscylacje, o czym pisałem w artykule *Wzmacniacz operacyjny, nie wzbudza*

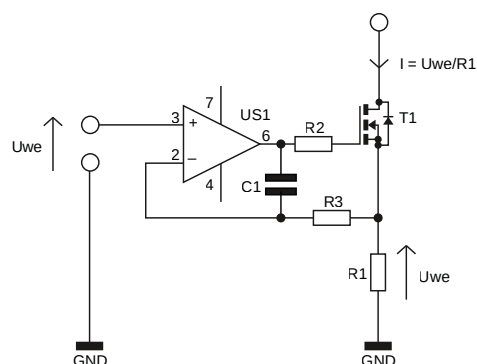


Rysunek 2. Zmniejszenie wpływu pojemności bramkowej tranzystora MOSFET na wyjście wzmacniacza operacyjnego

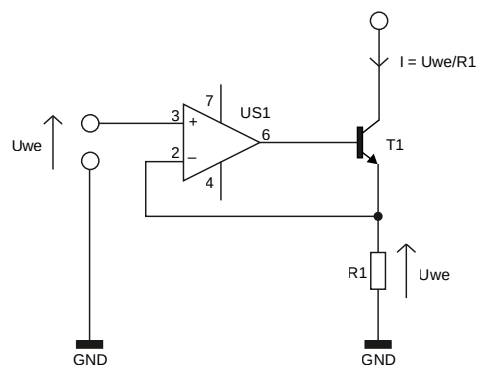
się!, opublikowanym w EP 09/2023. Mamy więc układ, który na papierze wygląda idealnie, ale w praktyce jest bardzo niesforny.

Co z tym fantem zrobić? Spora rzesza konstruktorów próbuje jakoś okiełznać tego potworka poprzez dodanie rezystora pomiędzy wyjściem wzmacniacza operacyjnego a bramką tranzystora – **rysunek 2**. Tak oto do układu doszedł rezystor R2. Jaką powinien mieć wartość? Niektórzy piszą, że 100 Ω wystarczy. Inni próbują wyprowadzać wzory na to. Niektórzy idą jeszcze dalej i zawężają pasmo przenoszenia całego tego układu kondensatorem C1 o niewielkiej pojemności (rzędu 100 pF) – **rysunek 3**.

Układ cały czas puchnie, a upragnionej stabilności nie widać, przynajmniej u mnie. Zrobiłem kiedyś prototyp układu zawierający źródło prądowe jak z rysunku 3 i doświadczalnie ustaliłem wartości elementów zapobiegających wzbudzeniu. Do próbnej serii 50 sztuk przyjąłem te wartości opatrzone kilkudziesięcioprocentowym zapasem. I co? Kilkanaście sztuk wzbudzało się bezlitośnie.



Rysunek 3. Przykład dążenia do poprawy stabilności układu



Rysunek 4. Tranzystor bipolarny nie powoduje wzbudzenia się układu

Pozostawało dalej kombinować z wartościami podzespołów, choć teraz na większą skalę.

Nie lubię układów, które działają w sposób magiczny. Niektórzy nie mają z tym problemów, dają zabezpieczenia przed wzbudzeniem (opracowane przez siebie już lata temu) i żyją szczęśliwi. Mam inne podejście: po co budować układy skłonne do wzbudzania się, skoro istnieje droga alternatywna – i to bardzo prosta? Oczywiście gdyby inaczej się nie dało – cóż, trzeba byłoby kombinować z poprawą stabilności istniejącego układu.

Mam na myśli zastąpienie tranzystora MOSFET bipolarnym – **rysunek 4**. Taki układ nigdy nikomu się nie wzbudził, chyba że elektronik wyjątkowo przesadził z niechlujnością montażu. Już załatwione, czas się rozejść, pora na CS-a? Ktoś może teraz wskoczyć przez okno, rozedrzeć koszulę na piersiach, niczym Rejtan na známym obrazie Jana Matejki – i krzyknąć: A CO Z PRĄDEM BAZY?!

Fakt, on istnieje i ma się dobrze. Wpływa też na nasze źródło prądowe, bowiem rzeczywisty prąd „wciągany” kolektorem jest mniejszy od tego, który przepływa potem przez emiter, a który wywołuje spadek napięcia na R1. Jesteśmy w stanie tak „pomanewrować” tym układem, aby skompensować ten wpływ odrażającego prądu bazy i zniwelować jego szkodliwe działanie niemal do zera. Pytanie, czy warto? Może nie trzeba się nim w ogóle przejmować?

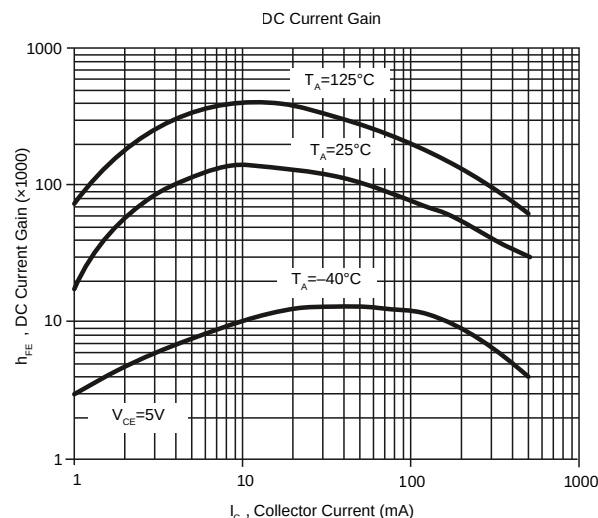
Zaraz, zaraz... Czyli dalibyśmy radę go pominąć, zaniedbać i żyć spokojnie dalej? Odpowiadam: w wielu przypadkach tak. Zastanówmy się: typowe wzmocnienie prądowe tranzystora małej mocy, jak BC846, zawiera się w przedziale od 200 do 450 (dla grupy wzmocnienia B) i od 420 do 800 dla grupy C (pomijam grupę A, bo jest już dzisiaj trudna do kupienia). Oznacza to, że prąd bazy będzie co najmniej 200 razy mniejszy od prądu kolektora. Co najmniej – bowiem podgrzanie tranzystorów powoduje wzrost wzmocnienia prądowego, marne mamy również nadzieje na znalezienie elementu najgorszego z danej grupy. Oznacza to, że tranzystor wprowadzi błąd ustalenia prądu wyjściowego o wartości 0,5% lub mniej. Cały czas trzeba mieć na uwadze, że do opisanego zjawiska dokładają się takie czynniki, jak wyżej wymienione, czyli:

- tolerancja rezystora R1 (1% lub mniej),
- offset napięciowy wzmacniacza operacyjnego (kilka miliwoltów lub mniej),
- potencjalnie nieskompensowany wpływ prądów wejściowych wzmacniacza operacyjnego.

Wartość R1 można wyregulować lub sięgnąć po rezystor z górnej półki, wzmacniacz operacyjny – wziąć z bardzo niskim offsetem napięciowym (choćby znany OP07), prądy wejściowe – prawidłowo skompensować lub zastosować wzmacniacz operacyjny z wejściami wykonanymi z użyciem tranzystorów unipolarnych. A ten nieszczęsny prąd bazy z nami pozostanie. Powiem więcej: będzie zależał od natężenia prądu wymuszanego przez źródło – oraz od temperatury i charakterystyki konkretnego egzemplarza elementu. Co więc z nim począć? Ja bym go zmniejszył.

Można sięgnąć po tranzystor z większym wzmocnieniem prądowym, ale takie są trudno dostępne. W zdecydowanej większości przypadków łatwiej, taniej, szybciej i prościej będzie sięgnąć po układ Darlingtona. Bardzo fajnym tranzystorem do tych zastosowań okazuje się BCV47, który napotkamy m.in. w obudowie SOT23 (ma układ wyprowadzeń taki sam jak BC846, jest niedrogi oraz ma wzmocnienie prądowe sięgające wielu – nawet kilkudziesięciu – tysięcy). Szczegóły obrazuje wykres zawarty w nocie katalogowej tego tranzystora, widoczny **rysunku 5**. Pomijając sytuacje ekstremalne, jak choćby silne mrozy i wysokie prądy (przy których i tak warto sięgnąć po inny element), możemy spodziewać się wzmocnienia rzędu 20 000 A/A.

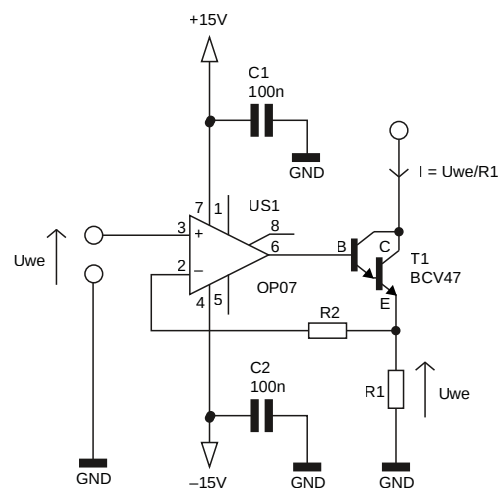
Konia z rzędem (albo w wersji nowoczesnej: SUV-a z garażem) temu, kto wychwyci błąd ustalenia prądu na poziomie 0,005% lub mniejszym. Okej, może w jakimś precyzyjnym sprzęcie pomiarowym niepewność tego rzędu byłaby problemem – dajmy na to, przy realizacji pomiaru rezystancji metodą czteropunktową – ale piszę



Rysunek 5. Zależność wzmocnienia prądowego tranzystora BCV47 od temperatury i prądu kolektora

tutaj o praktycznej realizacji układów przeznaczonych do seryjnej produkcji, w których redukcja kosztu każdego elementu przekłada się na wymierne korzyści dla producenta. Tam, gdzie w grę wchodzi indywidualna regulacja każdej wyprodukowanej sztuki, można stosować zupełnie inne podejście.

Reasumując: da się zbudować układ precyzyjnego źródła prądowego, które nie sprawi problemów z uruchomieniem, będzie cechowało się niską ceną podzespołów i utrzyma bardzo wysoką dokładność przetwarzania wejściowego napięcia na prąd. Proponowany przeze mnie układ można zobaczyć na **rysunku 6**. Elementem, którego dotychczas nie omawialiśmy, jest R2. Służy on do skompensowania rezystancji wewnętrznej źródła napięcia sterującego Uwe, aby oba wejścia wzmacniacza US1 były sterowane przez taką samą (lub choćby zbliżoną) rezystancję. Gdyby nie on, rezystancja sterująca wejściem odwracającym okazałaby się niemal równa zero. Warto więc zastosować R2 jak najbardziej zbliżony do rezystancji wewnętrznej źródła Uwe.



Rysunek 6. Przykład praktycznej realizacji źródła prądowego z tranzystorem Darlingtona

Tego typu układów powstały już setki, jeśli nie tysiące – i żaden z nich nie sprawiał problemów. Nawiasem mówiąc, wspomniana wcześniej seria prototypowa została „uratowana” przez wstawienie tranzystora BCV47, zamiast nieszczęsnego BSS123. Od tamtej pory nawet nie próbowałem wracać do idei precyzyjnego źródła prądowego z tranzystorem MOSFET sterowanym przez wzmacniacz operacyjny.

Michał Kurzela, EP



Nierozumiana pojemność kabli

Kable stosowane w technice audio do przesyłania sygnału analogowego mają budowę koncentryczną. Taki standard panuje od wielu lat, a jego niepodważalne zalety nie zwiastują zastąpienia tej technologii inną. Jednak z takim ukształtowaniem przewodu wiąże się niedogodność, na pierwszy rzut oka bolesna zwłaszcza w technice audio. Jaka?

Żyłą sygnałową, otuloną elastycznym dielektrykiem, a dookoła – metalowa warstwa przewodząca. Oto najprostszy opis kabla koncentrycznego, który od wielu lat służy do transmisji sygnałów o najrozmaitszym charakterze – od wrażliwych na zakłócenia informacji z analogowych czujników, poprzez sygnały wielkiej częstotliwości odebrane przez antenę telewizyjną, po fidera transmitujące wielką moc do anten nadawczych. Ich szczegółowa budowa różni się w zależności od zastosowania, ale ogólny kształt jest zawsze taki sam.

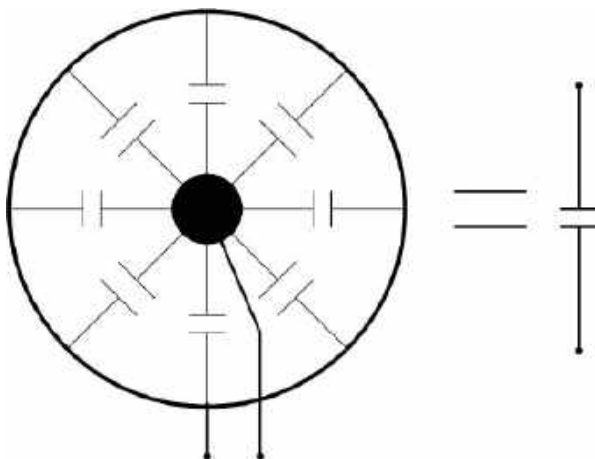
Ponieważ żyła sygnałowa jest „otulona” warstwą przewodzącą, tworzy się wokół niej klatka Faradaya o bardzo dobrych właściwościach ekranujących ją przed wpływem zewnętrznego pola elektromagnetycznego, głównie przed składową elektryczną. Tworzy to również wyśmienite warunki do propagacji fal wewnątrz takiego kabla, gdyż zewnętrzna warstwa ekranująca stanowi dobrą drogę powrotną dla prądu transmitowanego przez umieszczoną w środku żyłę.

Ale, ale, nie ma tak dobrze... Dwie warstwy przewodzące przedzielone najprzedejście dielektrykiem toż to przepis na pojemność!

I to nie byle jaką, bo wykazującą (z reguły) wysoką stałość w funkcji częstotliwości. Schematycznie pokazuje to **rysunek 1**. Odcinek kabla koncentrycznego możemy traktować jako pojemność, mimo że jest to podzespół o stałych rozłożonych, a nie skupionych. Nikt tam przecież na siłę kondensatorów nie wciska! Zaś skutki istnienia tej pojemności odczuć możemy jak najbardziej. W połączeniu z niezerową impedancją wyjściową stopnia zasilającego ów kabel tworzy ona bowiem filtr dolnoprzepustowy, w najprostszym wypadku jednobiegunowy – **rysunek 2**. To ogranicza pasmo przenoszenia i wydłuża odpowiedź impulsową, na szczęście jednak nie demoluje charakterystyki fazowej, bowiem opóźnienie wprowadzane przez taki człon RC jest stałe w szerokim przedziale częstotliwości.

Z jakimi wartościami liczbowymi możemy mieć do czynienia? Proszę bardzo, pierwszy z brzegu przewód koncentryczny TASKER C195 ma pojemność wynoszącą 130 pF/m. Jaką może mieć długość połączenie między lampowym przedwzmacniaczem a końcówką mocy? Przyjęcie wartości dwóch metrów (choć to nadmierne zapas) daje pojemność o wartości 260 pF. Celowo jako przykład wybrałem przedwzmacniacz lampowy, gdyż jego rezystancja wyjściowa może być znacznie wyższa niż w przypadku urządzeń tranzystorowych (choć nie jest to reguła obowiązująca w 100% sytuacji) – 20 kΩ w przypadku prostych przedwzmacniaczy, pozbawionych dodatkowych wtórników wyjściowych, nikogo mocno nie zdziwi.

Pojemność 260 pF i rezystancja 20 kΩ tworzą nam filtr dolnoprzepustowy o trzydecybelowej częstotliwości odcięcia wynoszącej nieco ponad 30 kHz. To nie jest jakaś kosmicznie wysoka wartość, oddziaływanie takiego filtra nachodzi już bowiem na pasmo



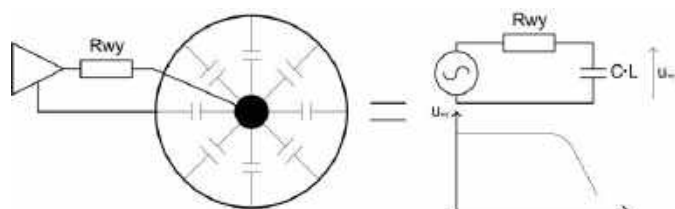
Rysunek 1. Uproszczona reprezentacja kabla koncentrycznego jako pojemności skupionej

akustyczne. Ciekawszy jest inny parametr, czyli stała czasowa takiego tworzywa RC, wynosząca $\tau=5,2 \mu\text{s}$. Będzie ona wpływała na czas narastania sygnału w torze audio, bowiem wydłuży go o $2,2\tau$, czyli do ponad $11 \mu\text{s}$. Niby niewiele, ale w systemach audio wysokiej jakości ten parametr podaje się równorzędnie z pasmem przenoszenia, dla przykładu: wzmacniacz Audio Research GSi75 ma ten parametr równy $40 \mu\text{s}$. W ostatnich latach coraz częściej mówi się o konieczności przejścia do analizy systemów audio w dziedzinie czasu, bo dziedzina częstotliwości nie oddaje wszystkich niuansów związanych z analizą brzmienia, ale to temat na oddzielny artykuł.

Mamy więc dodatkowy, niepożądany człon jednobiegunowy o czasie narastania wynoszącym $11 \mu\text{s}$ i częstotliwości odcięcia 30 kHz , co może dawać pewne tłumienie w zakresie wysokich częstotliwości. Dlatego warto dążyć do dwóch rzeczy: zmniejszenia pojemności kabla oraz zmniejszenia rezystancji, która nim steruje. Pierwszy parametr można w pewnym zakresie regulować, dobierając grubość i materiał dielektryka, przez co da się uzyskać kilkukrotną redukcję jego wartości. Z drugim parametrem jest jeszcze prościej, gdyż rezystancję wyjściową można kształtować w bardzo szerokim zakresie, zwłaszcza w przypadku urządzeń półprzewodnikowych. Nie należy jednak przesadzać w drugą stronę, gdyż może to doprowadzić do utraty stabilności, o czym pisałem w artykule „Wzmacniacz operacyjny, nie wzbudza się!”, opublikowanym w „Elektronice Praktycznej” 09/2023.

Przy znacznej długości kabla warto dążyć do uzyskania dopasowania. Poza najistotniejszym faktem, jakim jest dopasowanie impedancyjne i związana z nim transmisja możliwie dużej części energii, będzie ono miało również pozytywny wpływ na przytoczoną wyżej stabilność stopnia wyjściowego, które steruje niemal całą pojemnością. Może teraz inny przykład – źródło z wbudowanym przedwzmacniaczem (w celu zmniejszenia impedancji wyjściowej i uzyskania dopasowania), ale za to połączone kablem o znacznej długości, dajmy na to 80 m . Uzyskanie takiej długości w budynku czy na otwartej przestrzeni wcale nie jest czymś nadzwyczajnym.

Weźmy jakiś lepszej klasy przewód, proponuję 1703 SL005 od Alpha Wire. Jego pojemność to 36 pF/ft , czyli około 110 pF/m . Analizowany odcinek będzie zatem miał wypadkową pojemność $8,8 \text{ nF}$, a to już niemało – zakładając uproszczenie całej pojemności do jednego elementu skupionego. Jego impedancja charakterystyczna wynosi 43Ω (bardzo ładnie ze strony producenta, że podał ten parametr, nie zawsze jest on dostępny), więc chcąc uzyskać dopasowanie, trzeba sterować nim ze źródła o takiej właśnie rezystancji wyjściowej. Zatem pasmo przenoszenia takiego tworzywa wynosi teoretycznie 420 kHz .



Rysunek 2. Wpływ pojemności kabla na pasmo przenoszenia toru audio (C – pojemność kabla przeliczona na 1 m długości [F/m], L – długość przewodu [m])

Tyle że przy takim uproszczeniu sprawy nie byłoby możliwe transmitowanie kablami koncentrycznymi sygnałów radiowych, o telewizyjnych czy radarowych nawet nie wspominając. O rety, co to będzie, cała elektronika do kosza! Niekoniecznie cała, bo wyłącznie teoria obwodów – wygodna w użyciu i codziennej pracy nad układami, ale bezużyteczna w szerszym kontekście.

Przecież należy mieć na uwadze efekty falowe, a tak długiego kabla nie można już traktować jako pojedynczego elementu o stałych skupionych, bowiem nie spełnia on podstawowego, „teorioobwodowego” założenia: wymiary obwodu mają być mniejsze niż 10% długości fali w nim propagującej. Niektóre źródła podają 5% , spotkałem się też z wartością 20% , ale najczęściej trafia się wspomniane 10% . Zresztą to tylko pewne uproszczenie. Stała dielektryczna polietyleny, który jest izolatorem w tym kablu, wynosi $\epsilon_r=2,3$, zatem nie ma on przenikalności elektrycznej próżni. Fala elektromagnetyczna będzie zatem poruszała się w nim wolniej niż na otwartej przestrzeni.

Przekształcając kilka podstawowych wzorów dotyczących fal elektromagnetycznych, można uzyskać progową długość, powyżej której trzeba patrzeć na kabel koncentryczny jak na falowód:

$$L_{gran} = \frac{c}{10f \cdot \sqrt{\epsilon_r}} [m]$$

Można przyjąć, że we wszystkich kablach stosowany jest ten sam dielektryk (lub inny, ale o podobnej przenikalności elektrycznej), a szybkość światła w próżni wynosi $3 \cdot 10^8 \text{ m/s}$. Wtedy powyższy wzór upraszcza się do prostej zależności:

$$L_{gran} \approx \frac{20000000}{f} = \frac{2 \cdot 10^7}{f} [m]$$

lub po przekształceniu:

$$f = \frac{2 \cdot 10^7}{L_{gran}} [Hz]$$

Co z tego wynika? Kabel o długości 80 m możemy traktować „teorioobwodowo” dla częstotliwości do 250 kHz . O ile sygnał audio jeszcze mieści się w tym paśmie (choć to zależy, czy bierzemy pod uwagę takie czynniki jak czas narastania, czy tylko „surowe” pasmo przepustowe), o tyle jakkolwiek sygnał cyfrowy o zboczach bardziej stromych od Gubałówki będzie wymagał podejścia falowego. Jego harmoniczne już mogą w ten zakres wchodzić. W każdym jednak przypadku warto zadbać o dopasowanie nadajnika do kabla, aby zminimalizować odbicie energii od kabla ku nadajnikowi. Jeżeli zaś godzimy się na silne niedopasowanie (jak w przykładzie z przedwzmacniaczem lampowym), kiedy to impedancja wyjściowa źródła sygnału jest znacząco wyższa od impedancji charakterystycznej kabla, trzeba spojrzeć na jego pojemność. Może ona mieć niebagatelną rolę na jakość przeniesionego sygnału.

Michał Kurzela, EP

Bibliografia:

1. <https://t.ly/Fe3YB>
2. <https://t.ly/wsmv0>
3. <https://t.ly/EtjuH>



Konwersja analogowo-cyfrowa w technice audio (1)

W dziale „Audio bez tajemnic” nie mogło zabraknąć artykułów poświęconych fundamentalnym zagadnieniom związanym z cyfrowymi technikami rejestracji, przetwarzania i odtwarzania dźwięku. Na pierwszy ogień bierzemy konwersję analogowo-cyfrową – nieprzypadkowo zaczynamy właśnie od tej tematyki, gdyż bez nowoczesnych, szybkich i niskoszumowych przetworników ADC w ogóle nie byłibyśmy w stanie przejść od dźwięku w postaci sygnałów analogowych do świata cyfrowego.

Przyjęło się uważać, że wszystkie bodźce docierające do naszych zmysłów mają charakter analogowy. Oznacza to, że w każdym dowolnym momencie można określić ich konkretną wielkość. Dotyczy to również dwóch najważniejszych zmysłów: wzroku i słuchu. Nasz zmysł słuchu jest ciągle pobudzany zmianami ciśnienia otaczającego nas powietrza. Te zmiany są interpretowane przez nasze mózgi jako wrażenia dźwiękowe. Zmysł słuchu – bardzo istotny dla przetrwania – jest również źródłem wrażeń

estetycznych. Rytm i melodia towarzyszą nam od bardzo dawna. Kiedy rozwój nauki i techniki osiągnął odpowiednio wysoki poziom, ludzie zaczęli szukać sposobu na utrwalenie dźwięku i potem jego odtworzenie. Francuz Edward-Leon Scott de Martinville był pierwszym, który zapisał dźwięk na długo przed Edisonem. Jego fonautograf był zbudowany z tuby i membrany. Membrana, pobudzona falami dźwiękowymi, drgała i poruszała rylce dotykający papierowego cylindra. Edison wykorzystał jego prace i skonstruował słynny już fonograf, który zapisywał drgania rylca na aluminiowym walcu. Fonograf pozwalał na zapisywanie i odtwarzanie zarejestrowanego dźwięku. To był początek całej serii całkowicie mechanicznych urządzeń konsumenckich, które pozwalały na odtwarzanie nagrań. W końcowej fazie rozwoju dołączyły do tej grupy patefony z napędem sprężynowym, odtwarzające szybko obracające się płyty z szelaku. Dźwięk był zapisywany w postaci rowka na płaskiej powierzchni, a odtwarzany czysto mechanicznie przez zespół igły z membraną, połączoną z tubą. Płyty gramofonowe przeszły długą drogę rozwoju i nadal są używane z powodzeniem jako źródło wysokiej jakości dźwięku. Oczywiście nowoczesne gramofony nie są urządzeniami czysto mechanicznymi, ale sam proces

odczytu odbywa się dalej mechanicznie – tak jak w pierwszych fonografach.

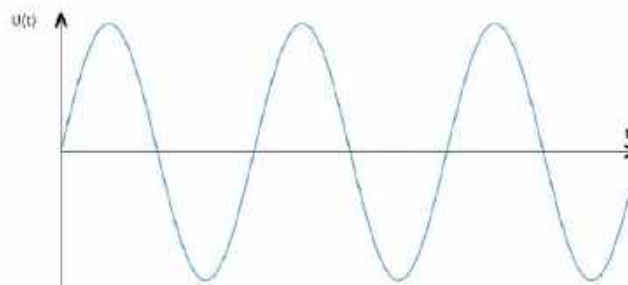
Drugą metodą utrwalania dźwięku jest zapis magnetyczny. W porównaniu z płytą gramofonową ma jedną niezaprzeczalną zaletę: możliwość nie tylko odtwarzania, ale też i wielokrotnego zapisywania materiału dźwiękowego, również w warunkach domowych. Zapis magnetyczny może zaoferować bardzo dobrą jakość dźwięku – w czasach intensywnego rozwoju magnetofonów powstały konstrukcje studyjne o doskonałych parametrach, przeznaczone do zapisywania materiału muzycznego, na podstawie którego tłoczono potem płyty gramofonowe. W obszarze zastosowań konsumenckich magnetofony zajmowały kluczową pozycję. Były obowiązkowym wyposażeniem melomana. Nawet na bardzo ograniczonym nośniku, jakim była kasetka Compact, potrafiąco uzyskać satysfakcjonującą jakość zapisywania i odtwarzania.

Jak już wspominałem, obie te metody dawały konsumentowi możliwość odtwarzania dźwięku na bardzo wysokim poziomie, oczywiście pod warunkiem posiadania urządzeń o odpowiedniej jakości. Obie też bazowały na naturalnym, analogowym sposobie działania. Czy miały jakieś wady? Podstawową była degradacja zapisanego sygnału podczas odtwarzania. Nawet najlepsze płyty gramofonowe ulegają mikrouszkodzeniom w czasie mechanicznego kontaktu igły wkładki z rowkiem płyty. Rozwój wkładek wymagających małych nacisków igły oraz ulepszenie technologii produkcji płyt ogranicza to zjawisko, ale nadal nie można go całkowicie wyeliminować. W magnetofonach następuje natomiast naturalne ścieranie warstwy magnetycznej w czasie jej kontaktu z głowicami i innymi elementami toru przesuwu. Ponadto taśma w trakcie pracy jest rozciągana i ulega częściowemu rozmagnesowaniu, nie można też nie wspomnieć o zużywaniu się samych głowic i mechaniki napędu taśmy.

Dla przeciętnego konsumenta wady te nie miały zbyt wielkiego znaczenia, ale zwolennicy wiernego odtwarzania muzyki borykali się z tym problemem. W latach 70. XX wieku pojawiła się technika, która w założeniu mogła dać bardzo wysoką jakość dźwięku i była, przynajmniej teoretycznie, odporna na degradację w czasie odtwarzania. Technika ta opierała się na konwersji analogowego sygnału audio na postać cyfrową, której materiał audio był zapisywany na płycie i odtwarzany promieniem lasera. Na końcu cyfrowe dane przekształcano na postać analogową. Był to standard płyty Compact Disc. Jeżeli płyta CD nie została uszkodzona mechanicznie, to jakość dźwięku nie zmieniała się, niezależnie od liczby odtworzeń. Żeby zachęcić konsumentów do zakupu nowych odtwarzaczy, użyto sloganów o „krystalicznie czystym dźwięku” i niesamowicie wysokiej jakości nagranej muzyki. Na początku rozwoju standardu CD nie było to prawdą i jakość materiału dźwiękowego ustępowała wyraźnie jakości płyt winylowych i magnetofonów wysokiej klasy. Wpływ na to miały głównie niedoskonałe przetworniki cyfrowo-analogowe. Ostatecznie jednak standard się przyjął i cyfrowa rewolucja stała się faktem, co miało kolosalne znaczenie dla całego rynku muzycznego. Dzisiaj, po wielu latach rozwoju, trudno sobie wyobrazić świat audio bez techniki cyfrowej. Zapewnia ona skalowaną jakość dostosowaną do wymagań konsumenta, bezproblemowe odtwarzanie, ale też zapisywanie materiału na dowolnym nośniku danych cyfrowych. W ostatnich latach kluczowa staje się możliwość przesyłania cyfrowego sygnału audio przez łącza internetowe. Kiedyś podstawowym źródłem sygnału był tuner radiowy, gramofon, magnetofon lub odtwarzacz CD. Teraz tę funkcję przejmuje streamer odtwarzający muzykę przesyłaną przez Internet lub zapisaną na lokalnych nośnikach.

Próbkowanie sygnału analogowego

Jak już wiemy, otaczający nas świat ma naturę analogową – jesteśmy do niego przystosowani ewolucyjnie i dlatego dźwięki do nas docierające muszą być analogowe, co oznacza, że w dowolnym



Rysunek 1. Analogowy, sinusoidalny sygnał elektryczny

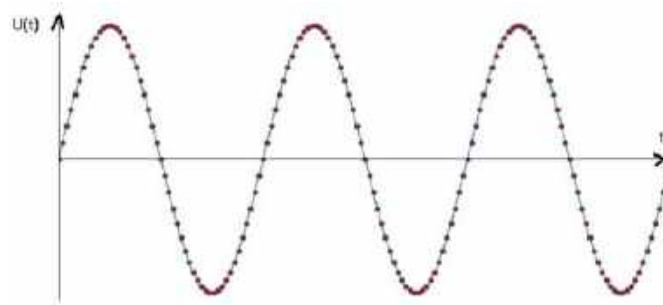
wybranym momencie można określić wartość poziomu natężenia dźwięku, a – po przekształceniu na sygnał elektryczny – także wartość napięcia odpowiadającego temu natężeniu. Jeżeli na wykresie zaprezentujemy czas w osi X, a wartość napięcia w osi Y, to możemy graficznie pokazać zmienność w czasie takiego sygnału – **rysunek 1**.

Zamiana sygnału analogowego na postać cyfrową nazywa się próbkowaniem i kwantyzacją – polega na mierzeniu analogowej z określoną częstotliwością oraz zapisywaniu jej wartości dyskretyzowanej. Po takiej operacji analogowy sygnał jest reprezentowany przez zbiór liczb – mówimy teraz, że sygnał ma reprezentację dyskretną. Ma to fundamentalne znaczenie, bo same liczby można bez problemu zapisywać w plikach, kopiować bez utraty jakości lub przysyłać dowolnymi łączami cyfrowymi, w tym przez Internet. A to nie koniec, bo możliwa jest kompresja tych danych – połączona z utratą jakości, ale powodująca, że ilość potrzebnych danych do zapisania utworu może drastycznie zmaleć, przy wciąż jeszcze akceptowanej jakości. Cyfrowa reprezentacja umożliwia ponadto matematyczne filtrowanie (funkcja *equalizer*), regulację poziomu sygnału, miksowanie itp.

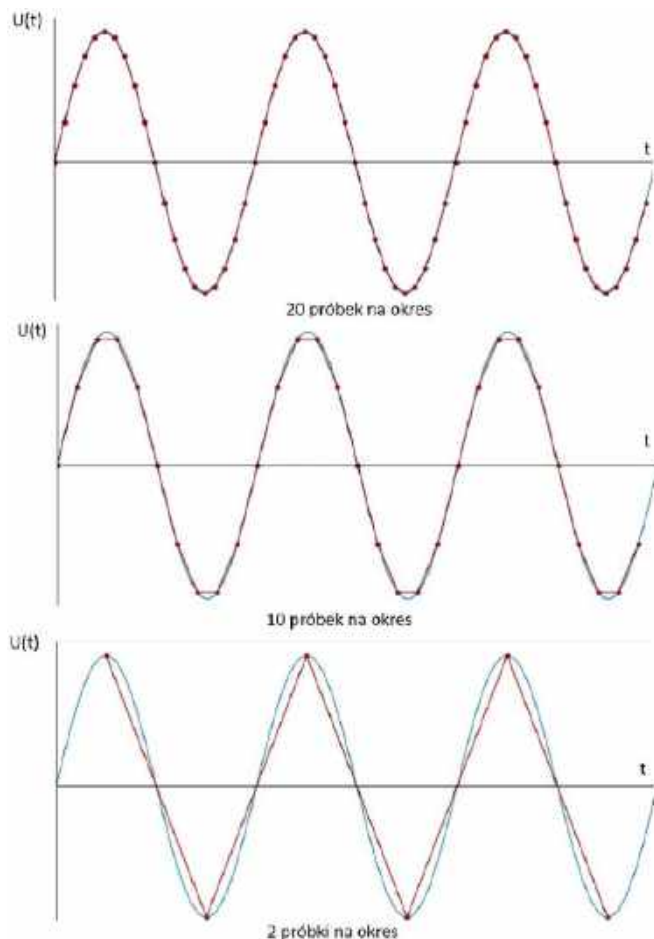
Kwantowanie wydaje się prostą operacją, ale niestety taka nie jest. Po pierwsze musimy sobie zadać pytanie: z jaką częstotliwością trzeba próbować (pobierać próbki), żeby potem z wynikowych danych móc odtworzyć sygnał analogowy bez utraty jego jakości? Intuicyjnie czujemy, że powinno się to odbywać z jak najwyższą częstotliwością. Im więcej próbek, tym wierniej będą reprezentowały sygnał próbkowany.

Na **rysunku 2** pokazano próbkowanie sygnału sinusoidalnego. Liczba próbek jest tak dobrana, że można z nich w przybliżeniu odtworzyć oryginalny sygnał. Im więcej będzie takich próbek, tym lepsze odwzorowanie. Takie postępowanie ma jednak sporą wadę: po pierwsze rośnie nam bardzo mocno ilość danych przeznaczonych do reprezentowania sygnału. Po drugie – nawet przy dużej liczbie próbek nie są one w stanie idealnie odtworzyć próbkowanego sygnału. Czyżby zatem cyfryzacja zapisu, która miała zapewnić bardzo dobrą jakość, z zasady powodowała zawsze zmianę kształtu przebiegu, a co za tym idzie – wprowadzała nieuchronne zniekształcenia harmoniczne? Na szczęście nie jest aż tak źle, ale powiemy o tym później.

Wróćmy do pytania o minimalną częstotliwość próbkowania. Wiemy, że jej zwiększanie nie spowoduje, że próbki będą reprezentowały ciągle sygnał. To oczywiste, bo próbkowanie zmienia



Rysunek 2. Próbkowanie sygnału sinusoidalnego



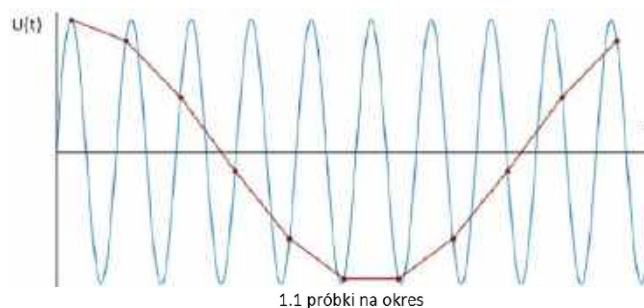
Rysunek 3. Próbkowanie tego samego przebiegu z częstotliwościami 20, 10 oraz 2 próbek na okres

naturę sygnału z ciągłego w dyskretny, to znaczy taki, w którym informacja o sygnale jest dostępna tylko w określonych interwałach czasowych. Pomiędzy tymi interwałami nie wiemy nic o jego kształcie. Dodajmy na marginesie, że innym znanym, dyskretnym przekształceniem sygnału ciągłego jest... przekształcenie Fouriera, w którym sygnał jest reprezentowany przez dyskretny poziom amplitudy sygnału o częstotliwości podstawowej oraz jej harmonicznych.

Jako przykład weźmy sygnał sinusoidalny o częstotliwości f . Wykonajmy próbkowanie równomierne tego sygnału z różnymi częstotliwościami: 20 próbek na okres, 10 próbek na okres i 2 próbki na okres (**rysunek 3**). Wraz ze zmniejszeniem częstotliwości próbkowania kształt przebiegu (po interpolacji liniowej) coraz mniej przypomina oryginalną sinusoidę. Ale w każdym z tych przypadków można jeszcze zobaczyć okresowy charakter sygnału próbkowanego i określić jego częstotliwość. Nawet wtedy, gdy będą to tylko dwie próbki na okres.

A co się stanie, kiedy zmniejszymy liczbę próbek poniżej dwóch na okres? Na **rysunku 4** pokazano próbkowanie z częstotliwością 1,1-f.

Nie dość, że nie możemy określić teraz częstotliwości próbkowanego sygnału, to proces próbkowania spowodował pojawienie się sygnału o znacznie niższej częstotliwości, której... nie było w oryginalnym przebiegu. Z analizy rysunków 3 i 4 możemy wysnuć luźne przypuszczenie, że częstotliwość próbkowania nie może być mniejsza niż 2 próbki na okres, gdyż tylko po spełnieniu tego warunku można określić częstotliwość oryginalnego zapisu. Poniżej tego progu próbki mogą reprezentować sygnały o zafalszowanych częstotliwościach, których nie ma w oryginalnym sygnale – co rzecz jasna nie powinno mieć miejsca.



Rysunek 4. Próbkowanie z częstotliwością 1,1 próbki na okres

W ten sposób dochodzimy do fundamentalnego, powszechnie znanego twierdzenia o próbkowaniu, które mówi, że **częstotliwość próbkowania f_s sygnału analogowego musi być przynajmniej dwukrotnie większa niż częstotliwość f_0 sygnału próbkowanego**. Możliwe jest wtedy idealne odtworzenie oryginalnego sygnału z zapisanych próbek. Jeżeli częstotliwość próbkowania jest niższa, to taka operacja okazuje się niemożliwa do przeprowadzenia.

Twierdzenie o próbkowaniu było możliwe do sformułowania dzięki pracom kilku uczonych, głównie Harry'ego Nyquista, Claude'a Shannona, ale też Edmunda Whittakera czy Władimira Kotelnikowa.

Harry Nyquist jest być może znany Czytelnikom jako twórca teorii warunków stabilności wzmacniaczy ze sprzężeniem zwrotnym, ale nazwisko uczonego chyba znacznie częściej kojarzy się z kryterium Nyquista, dotyczącym właśnie minimalnej częstotliwości próbkowania sygnałów analogowych, przy której da się z próbek odtworzyć oryginalny sygnał bez zniekształceń.

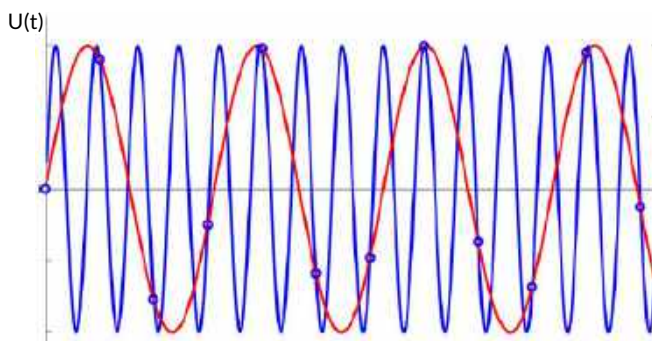
Niejednoznaczność sygnału w dziedzinie częstotliwości

Dyskretnie próbkowanie ciągłego sygnału ma pewną bardzo ważną właściwość. Załóżmy, że będziemy próbkować sygnał sinusoidalny z częstotliwością 4 próbki na okres. Takie próbkowanie spełnia warunek Nyquista, więc na podstawie skwantowanego zapisu możemy wykreślić oryginalny przebieg. Jednak kiedy poprosimy o to kogoś, kto nie zna oryginalnego przebiegu i nic o nim nie wie, to na podstawie tych próbek może wykreślić zupełnie inny przebieg. Taka możliwa sytuacja została pokazana na **rysunku 5**.

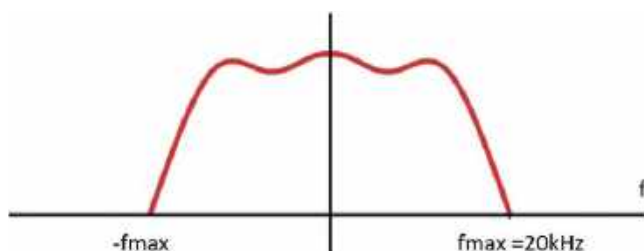
Jeżeli ciąg danych reprezentuje próbki przebiegu sinusoidalnego, to tylko na ich podstawie i bez dodatkowych informacji nie możemy określić jednoznacznie częstotliwości tego przebiegu.

Na **rysunku 5** widać – wyznaczone na podstawie dostępnych próbek – dwa przebiegi sinusoidalne. Zadajmy sobie pytanie, czy faktycznie da się wykreślić tylko dwa? Otóż okazuje się, że można takich sygnałów wskazać... nieskończenie wiele!

Założmy, że przebieg sinusoidalny o częstotliwości f_0 próbkujemy z częstotliwością f_s . Jeżeli k jest dowolną liczbą całkowitą, **to nie**



Rysunek 5. Niejednoznaczność sygnału w dziedzinie częstotliwości



Rysunek 6. Ciągłe widmo rzeczywistego, dolnopasmowego sygnału audio

jesteśmy w stanie odróżnić (bazując tylko na podstawie próbek) przebiegu o częstotliwości f_0 oraz sygnałów o sinusoidalnych o częstotliwościach $f_0 + k \cdot f_s$, gdzie k – dowolna liczba całkowita.

Co to oznacza w praktyce? Okazuje się, że nie istnieje taki ciąg danych, który mógłby reprezentować jedną i tylko jedną sinusoidę. Żeby z takiego ciągu odtworzyć próbkowany przebieg, potrzebne są dodatkowe informacje. Z tego również wynika, że proces próbkowania „powieli” częstotliwość f_0 do nieskończoności.

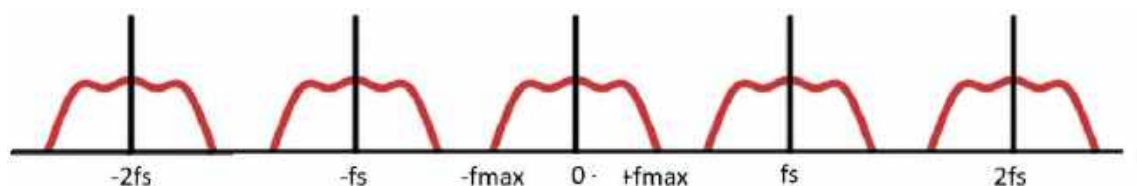
Powielanie widma, aliasing

Do tej pory ograniczyliśmy – dla wygody – nasze rozważania tylko do próbkowania pojedynczego sygnału sinusoidalnego o określonej częstotliwości. Rzeczywisty, ciągły sygnał audio jest jednak mieszaną różnych częstotliwości o charakterze dolnopasmowym. Przyjęło się, że widmo sygnału audio rozciąga się od 20 Hz do 20 kHz, ale my będziemy dla uproszczenia rozważać pasmo od 0 Hz do 20 kHz.

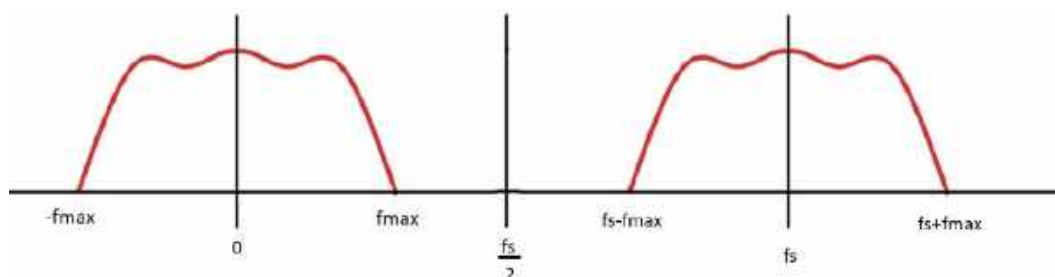
Analizowanie sygnałów audio w dziedzinie czasu jest dobre dla pojedynczego przebiegu sinusoidalnego, ale dla ciągłych sygnałów o ograniczonym paśmie dalsze rozważania wygodniej będzie przeprowadzać w dziedzinie częstotliwości. Żeby wyznaczyć widmo sygnału audio o charakterze dolnoprzepustowym, poddamy go przekształceniu Fouriera. Wiemy już, że częstotliwość graniczna wynosi $f_{\max} = 20$ kHz. Ciągłe widmo takiego sygnału zostało pokazane na **rysunku 6**.

Cała energia sygnału mieści się w zakresie od $-f_{\max}$ do $+f_{\max}$, a poza tym zakresem jest równa zero.

Po spróbkowaniu sygnału o takim widmie otrzymamy dyskretnie widmo sygnału próbkowanego o szerokości od $-f_{\max}$ do $+f_{\max}$ oraz szereg powielonych widm dyskretnych o takiej samej szerokości. Każde z nich jest przesunięte na osi częstotliwości o wartość $k \cdot f_s$. Zostało to pokazane na **rysunku 7**.



Rysunek 7. Powielanie dyskretnego widma po próbkowaniu sygnału o widmie z rysunku 6

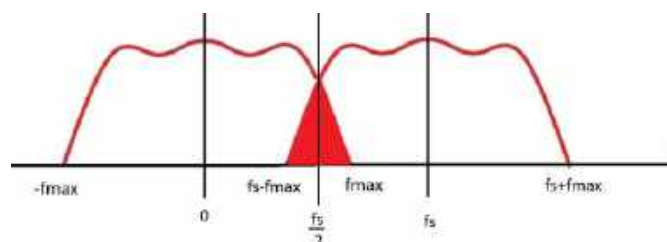


Rysunek 8. Powielanie widma przy spełnionym warunku $f_s > 2 \cdot f_{\max}$

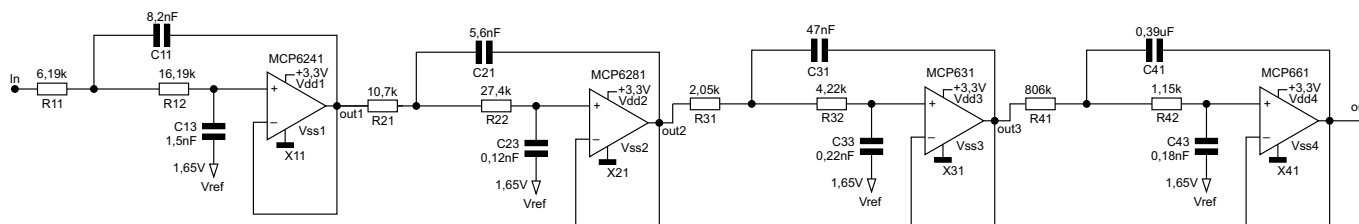
Możemy się zastanawiać, dlaczego te powielane widma stanowią potencjalne źródło problemów? Przecież interesują nas częstotliwości od $-f_{\max}$ do $+f_{\max}$, a – jak widać na rysunku 7 – to widmo nie jest w żaden sposób modyfikowane. Żeby tak było w istocie, trzeba znać i prawidłowo stosować sposoby kontrolowania położenia powielanych widm na osi częstotliwości, właśnie za pomocą odpowiedniego doboru częstotliwości próbkowania. Jak się łatwo domyślić, kluczowym aspektem będzie kryterium Nyquista: $f_s > 2 \cdot f_{\max}$. Jeżeli jest ono spełnione, to powielane widma nie zachodzą na widmo podstawowe, bo $f_{\max} < f_s/2$ – tak, jak to zostało pokazane na **rysunku 8**. Warunek niezachodzenia na siebie powielonego widma i widma sygnału próbkowanego jest koniecznym warunkiem do idealnego odtworzenia sygnału analogowego z próbek, które utworzyliśmy w procesie próbkowania. Powiemy o tym jeszcze za chwilę.

Zobaczmy, jak będzie wyglądała sytuacja, gdy kryterium Nyquista nie zostanie spełnione – spójrzmy na **rysunek 9**. Widmo powielone zachodzi na widmo sygnału próbkowanego, powodując w rezultacie nieodwracalne zniekształcenia widma sygnału oryginalnego. To zjawisko nazywa się aliasingiem. Jeżeli w trakcie próbkowania powstanie aliasing, to nie mamy możliwości, żeby odtworzyć w przetworniku cyfrowo-analogowym sygnał oryginalny na podstawie zapisanych próbek. Jak już wspomniałem, opisywane zjawisko niszczy nieodwracalnie cyfrową reprezentację oryginalnego sygnału i musi być koniecznie wyeliminowane.

Podstawowy środek zapobiegawczy jest nam znany – wystarczy zastosować w praktyce kryterium Nyquista. Do sygnałów audio w paśmie ograniczonym do 20 kHz stosuje się dwie standardowe częstotliwości próbkowania: 44,1 kHz i 48 kHz. Teoretycznie powinno to wyeliminować problem aliasingu. Jednak w rzeczywistych układach nie istnieją idealne sygnały dolnoprzepustowe, w których powyżej $f_{\max}$ energia sygnału jest zerowa. W torach przesyłowych jest zawsze obecny szerokopasmowy szum. Częstotliwości



Rysunek 9. Powielanie widma przy niespełnionym warunku $f_s < 2 \cdot f_{\max}$



Rysunek 10. Schemat analogowego filtra antyaliasingowego

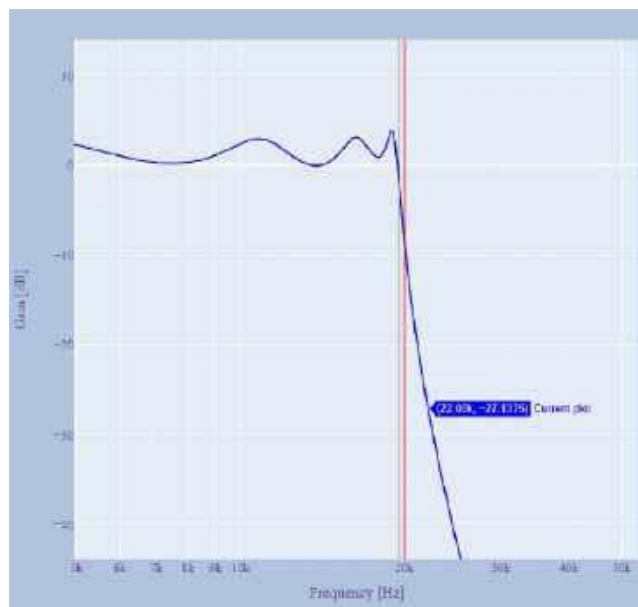
tego szumu znacznie wykraczają poza częstotliwości graniczne $f_{\max}$. Poza tym istnieje szereg zakłóceń radiowych bądź przewodzonych (z sieci energetycznej), które mogą być obecne w sygnale analogowym. Dlatego na wejściach przetworników analogowo-cyfrowych stosuje się dolnoprzepustowe filtry antyaliasingowe. Mają one za zadanie odfiltrowanie częstotliwości powyżej $f_{\max}$.

Filtry antyaliasingowe

Analogowe filtrowanie antyaliasingowe jest dość trudne w realizacji. Jeżeli przyjmijmy, że $f_{\max}$ jest równa 20 kHz, a f_s ma wartość 44,1 kHz, to filtr powinien przenosić bez zniekształceń częstotliwość 20 kHz i całkowicie tłumić częstotliwość 22,05 kHz ($f_s/2$). Wykonanie takich filtrów jest technicznie kłopotliwe, a na pewno skomplikowane i drogie. Na stronie firmy Microchip można znaleźć kreator filtrów analogowych FilterLab pracujący on-line i wykonać proste ćwiczenie symulacji projektu filtra dolnoprzepustowego, który będzie przenosił pasmo 20 kHz ze spadkiem -3dB . Wybrałem filtr Czebyszewa 8. rzędu w topologii Sallena-Keya. Schemat filtra został pokazany na rysunku 10, a jego charakterystyka amplitudowa – na rysunku 11.

Filtr Czebyszewa charakteryzuje się szybkim wzrostem tłumienia w pasmie zaporowym, ale niestety wprowadza zafalowania charakterystyki w pasmie przepustowym. Jeżeli będziemy próbować sygnał o szerokości widma 20 kHz z częstotliwością $f_s = 44,1\text{ kHz}$, to przy częstotliwości $f_s/2$ tłumienie naszego filtra będzie równe ok. -27 dB . Jeżeli próbkowany sygnał będzie miał duży poziom zakłóceń powyżej 20 kHz, to nawet taki filtr nie zapewni odpowiedniego tłumienia zjawiska aliasingu.

Na pierwszy rzut oka wydawałoby się, że może w dużych studiach nagraniowych nie stanowiłoby problemu zastosowanie skomplikowanych i rozbudowanych filtrów o odpowiednich parametrach, zapewniających bardzo niski poziom aliasingu. Jednak znaleziono dużo lepszy sposób na pozbycie się aliasingu niż rozbudowane filtry analogowe na wejściu przetworników analogowo cyfrowych.

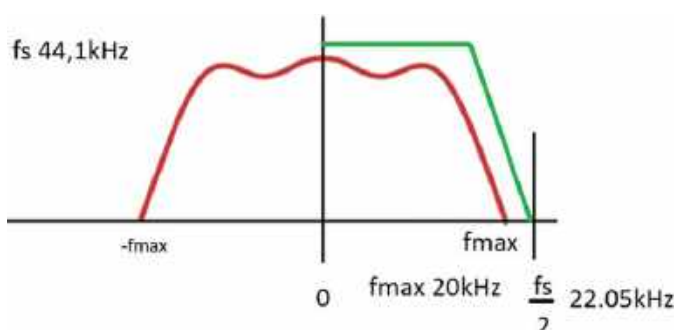


Rysunek 11. Charakterystyka amplitudowa filtra z rysunku 10

Nadpróbkowanie (oversampling) w konwersji analogowo cyfrowej

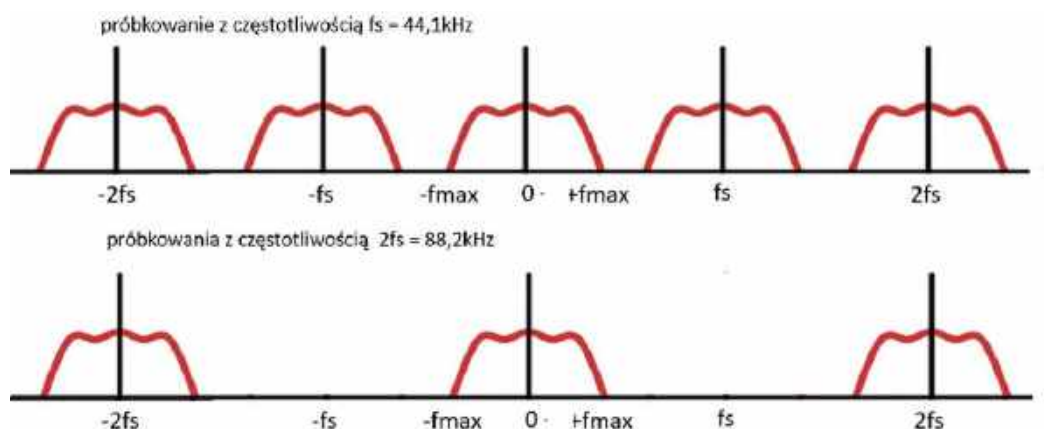
Na rysunku 12 pokazano typową sytuację spotykaną w technice audio. Sygnał analogowy o pasmie 20 kHz jest próbkowany z częstotliwością 44,1 kHz. Żeby wyeliminować aliasing, w sygnale analogowym muszą być odfiltrowane wszystkie częstotliwości powyżej $f_s/2$, czyli 22,05 kHz. Na wykresie pasma zaznaczono wymaganą charakterystykę filtra dolnoprzepustowego.

Jak już wiemy, implementacja takiego filtra jest trudna pod względem technicznym. Kryterium Nyquista nakłada na nas konieczność próbkowania z częstotliwością nie mniejszą niż $2 \cdot f_{\max}$ – możemy więc z powodzeniem próbować z częstotliwościami większymi. Zobaczmy co się stanie, gdy będziemy rejestrować sygnał z częstotliwością dwa razy większą, niż wynikająca z kryterium Nyquista. Powielone pasma na osi częstotliwości będą rozmieszczone nie co wartości f_s , ale co $2 \cdot f_s$ – tak jak to zostało pokazane na rysunku 13.

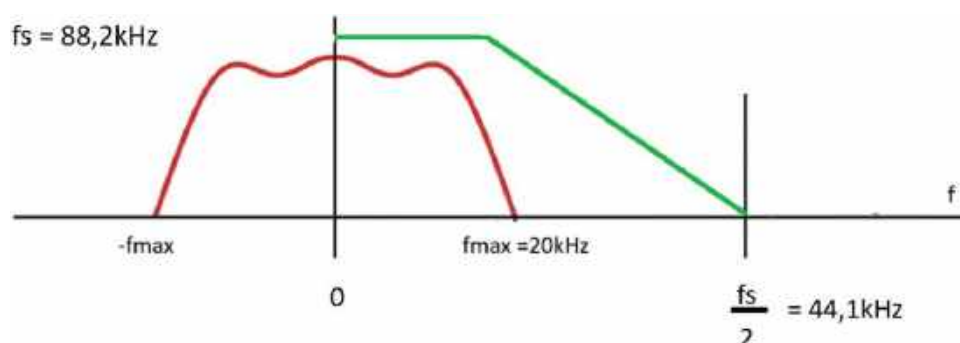


Rysunek 12. Filtrowanie sygnału analogowego przy próbkowaniu sygnału audio z częstotliwością 44,1 kHz





Rysunek 13. Nadpróbkowanie sygnału audio



Rysunek 14. Filtrowanie sygnału analogowego przy dwukrotnym nadpróbkowaniu sygnału audio z częstotliwością 88,2 kHz

Nadpróbkowanie to proces polegający na zwielokrotnianiu częstotliwości próbkowania wynikającej z kryterium Nyquista. W naszym przykładzie pasmo analogowe jest nadal ograniczone do 20 kHz, ale próbkujemy je z częstotliwością nie 44,1 kHz, ale 88,2 kHz. W efekcie takiej operacji pierwsze powielone pasmo zaczyna się dużo dalej od końca pasma podstawowego. Co nam to daje? Filtr dolnopasmowy musi tłumić nie do 22,05 kHz, tylko do $f_s/2$, czyli do 44,1 kHz. Taki filtr jest już dużo łatwiejszy do wykonania.

No dobrze, nadpróbkowanie eliminuje paskudny problem ze skomplikowanym filtrem antyaliasingowym, ale generuje kolejny kłopot. Jeżeli próbkujemy sygnał z częstotliwością wielokrotnie większą niż wymagana do idealnego odtworzenia sygnału z próbek, to mamy mnóstwo nadmiarowych próbek, które w istocie... nie są do niczego potrzebne. W praktyce przydałyby się one w trakcie konwersji cyfrowo-analogowej, bo tam będzie lustrzany problem z powielaniem pasm, tylko już w domenie analogowej, a nie w dyskretnej. Ale konwersja C/A poradzi sobie i bez tego, a nadmiarowa ilość danych z konwersji analogowo-cyfrowej bardzo komplikuje ich zapisywanie i przesyłanie.

Decymacja (downsampling)

Technika cyfrowa radzi sobie z tym problemem za pomocą operacji decymacji. Decymacja to zachowanie co n -tej próbki i odrzucenie pozostałych ze zbioru nadpróbkowanych danych – bez szkody dla informacji, z których będziemy odtwarzać później sygnał analogowy. Przykładowo: w przypadku dwukrotnego nadpróbkowania zachowujemy jedynie co drugą próbkę.

Zastanówmy się jednak nad pewną kwestią. Dwukrotnie nadpróbkowane dane reprezentują sygnał o dwa razy szerszym pasmie: w naszym przypadku nie 20 kHz, a 40 kHz (a dokładniej 44,1 kHz). Wykonanie decymacji będzie znowu przesunęło powielone widmo w kierunku niższych częstotliwości, co powoduje

ponownym pojawieniem się aliasingu, którego przecież chcieliśmy uniknąć zwiększając częstotliwość próbkowania. Żeby poradzić sobie z tym problemem, trzeba najpierw dane przefiltrować dolnoprzepustowym filtrem cyfrowym (aby próbki ponownie reprezentowały sygnał o paśmie pierwotnego sygnału analogowego). Dopiero po tym poddajemy je decymacji. Jednak te operacje, w odróżnieniu do rozbudowanych filtrów analogowych, można wykonać bez problemu w scalonych przetwornikach analogowo-cyfrowych – już za miesiąc przyjrzymy bliżej temu zagadnieniu.

Tomasz Jabłoński, EP





Małe, ale wariaty: nowoczesne, scalone wzmacniacze audio małej mocy

Wzmacniacz audio... to brzmi dumnie!

Prawdopodobnie większość z nas pojęcie wzmacniacza kojarzy w pierwszej chwili ze sporym i ciężkim sprzętem, zdolnym do napędzania głośników o mocy kilkudziesięciu czy kilkuset watów. A przecież wokół aż roi się od „małych bohaterów” naszej konsumenckiej codzienności. Miniaturowe wzmacniacze nosimy w plecakach, torbach, kieszeniach, a także na uszach oraz wewnątrz nich. I właśnie tym niepozornym układom scalonym przyjrzymy się w niniejszym przeglądzie wybranych propozycji z ofert producentów półprzewodników. Przy okazji przemycę szereg ważnych trików układowych i aspektów konstrukcyjnych charakterystycznych głównie dla wzmacniaczy słuchawkowych. Zaczynamy!

Dawniej wszystkie wzmacniacze (podobnie jak i wszelkie inne urządzenia elektroniczne) były – jakże by inaczej – budowane z elementów dyskretnych. Wprowadzenie tranzystorów szybko i niemal całkowicie wyparło z rynku lampy elektronowe i to na dobre kilkadziesiąt lat. Pocziwe triody czy pentody wróciły do łask dopiero za sprawą sentymentu audiofilów i miłośników brzmienia (oraz wyglądu) retro. Przejście na konstrukcje scalone dokonało kolejnej rewolucji, a objętość i masa urządzeń znów drastycznie zmalały. Dziś wzmacniacze (głównie te pracujące w klasie D, choć nie tylko) mają nierzadko wymiary niewiele większe od „kosteczki” o rozmiarach milimetra i – jakby tego było mało – oferują szereg ciekawych funkcjonalności. Jakich? Odpowiedź na to pytanie znajdziecie, drodzy Czytelnicy, w tym właśnie artykule. Na potrzeby inauguracji

nowego działu dotyczącego audio przygotowałem zestawienie wybranych układów od takich gigantów, jak Texas Instruments, ST Microelectronics, Analog Devices, Maxim (dziś już także pod banderą Analoga) czy ams OSRAM.

Przegląd jest subiektywny, ale reprezentatywny – starałem się pokazać jak najszersze spektrum możliwości wzmacniaczy scalonych dostępnych na rynku. Gwoli ścisłości dodam, że skupiłem się głównie na wzmacniaczach małej mocy, tj. od kilkudziesięciu miliwatów do kilku watów. Wśród docelowych aplikacji takich maluchów można wskazać m.in.:

- słuchawki bezprzewodowe (nauszne, douszne, dokanałowe),
- niewielkie głośniki Bluetooth,
- smartfony,
- smartwatche,
- tablety,
- zestawy głośnomówiące,
- radioodbiorniki przenośne,
- nawigacje samochodowe,
- mobilne odtwarzacze MP3,
- aparaty i kamery cyfrowe z funkcją odtwarzania wideo,
- dyktafony i przenośne rejestratory audio,
- tłumaczniki mobilne,
- instrumenty muzyczne,
- urządzenia do gier,
- asystentów głosowych,
- systemy smart home,
- roboty sprząające,
- zabawki elektroniczne,
- monitory komputerowe z wbudowanymi głośnikami,
- urządzenia medyczne z funkcją odtwarzania komunikatów i instrukcji głosowych,

- wagi z funkcją głosową,
- elektroniczne dzwonki do drzwi,
- infokioski,
- interaktywne systemy multimedialne,

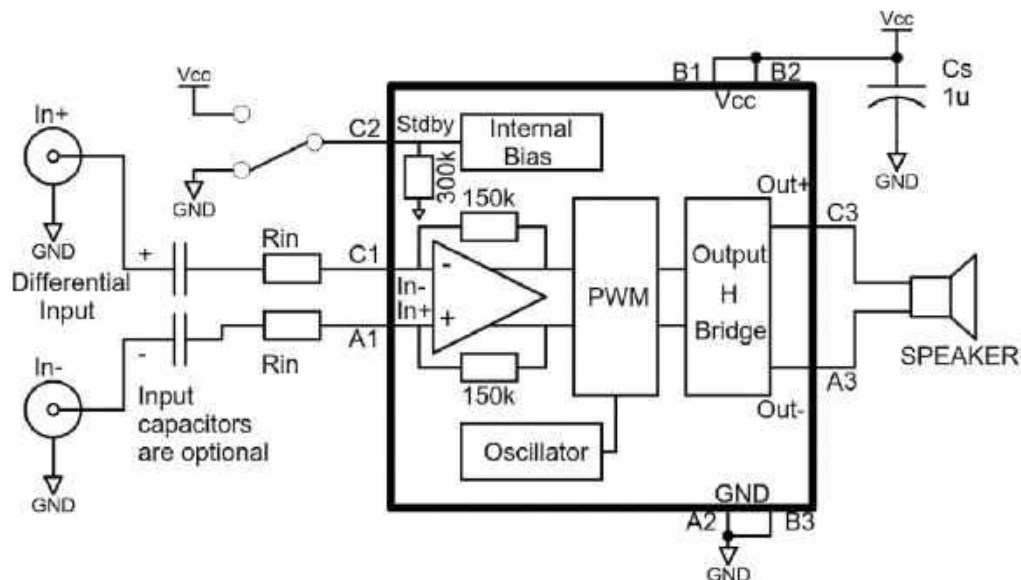
i wiele, wiele innych. Słowem – wszystkie te produkty, w których trzeba zaimplementować generowanie komunikatów akustycznych lub odtwarzanie plików audio, a zatem wytwarzać jakiegokolwiek dźwięku bardziej złożone, niż proste sygnały buzzera elektromagnetycznego lub piezoelektrycznego. A skąd akurat taki wybór? Powodów jest kilka.

Po pierwsze: ta specyficzna grupa komponentów jest przeznaczona przede wszystkim do urządzeń kompaktowych. Liczą się więc nie tylko wymiary samego układu, ale też prostota implementacji i możliwie krótki BOM. Miniaturyzacja wywiera szczególnie mocną presję na producentów, którzy – by przebić się w tym ciekawym obszarze rynku – muszą jak najwięcej funkcji „upchać” na płytce krzemowej. Mało tego – pożądaną cechą wzmacniaczy jest możliwość pracy bez wyjściowych filtrów dolnoprzepustowych, które w przypadku klasy D wydawać się mogą absolutnym *must have*.

Po drugie, aplikacje ubieralne i przenośne mają pewne szczególne potrzeby funkcjonalne. Przykład? Spora część dostępnych na rynku słuchawek ma funkcję aktywnego tłumienia szumów (ANC). Najogólniej rzecz biorąc działa ona poprzez „dorzucanie” do sygnału audio „kopii” dźwięku dobiegającego do uszu słuchacza z otoczenia, ale w przeciwfazie. W ten (tylko pozornie!) prosty sposób zaprzęgamy fizykę, a ściślej rzecz ujmując – interferencję fal dźwiękowych, w szczytnym celu poprawy wrażen odsłuchowych. Implementacja ANC ma spory sens w ograniczonej przestrzeni kanałów słuchowych, ale byłaby niemożliwie trudna – a w praktyce niemożliwa do wdrożenia – np. w całym pomieszczeniu, w którym dźwięk do słuchacza dobiega różnymi drogami (także poprzez wielokrotne odbicia). Dlatego na rynku pojawiły się nieliczne układy scalone udostępniające konstruktorom gotowy do użycia „bloczek” ANC.

Aktywna redukcja szumów jest zresztą przykładem „dużego kalibru” – znacznie bardziej przyziemna, choć także ważna z użytkowego punktu widzenia, jest opcja obsługi trybu bypass, dostępna w niektórych układach wzmacniaczy. Jej działanie polega na umożliwieniu odtwarzania dźwięku za pomocą nausznych słuchawek Bluetooth, których zasilanie jest wyłączone. Taka sytuacja może mieć miejsce np. po rozładowaniu akumulatora. Udostępnienie użytkownikom gniazda wejścia liniowego pozwala „napędzać” przetworniki bezpośrednio (pasywnie), z pominięciem układu elektronicznego. I choć funkcje ANC, regulacja głośności oraz inne cuda techniki audio pozostają niedostępne przy braku zasilania, to fakt, że w ogóle da się w takich warunkach używać ulubionych słuchawek, niejednemu miłośnikowi muzyki uratował już życie w sytuacji podbramkowej. Tak – piszę to m.in. na podstawie własnego doświadczenia.

Jeszcze inna opcja to wzmacniacz integrujący w sobie stereofoniczne wyjście słuchawkowe oraz dodatkowy kanał do obsługi



Rysunek 1. Podstawowy schemat aplikacyjny układu A21SP16 [1]

niewielkiego głośnika. Takie rozwiązanie może okazać się przydatne np. w projekcie przenośnego radioodbiornika.

Przykładów różnic jakościowych pomiędzy klasycznymi końcówkami mocy, a małymi wzmacniaczami słuchawkowymi i podobnymi można by wskazać jeszcze więcej. Nie przedłużajmy jednak wstępu i zanurzymy się w świat półprzewodnikowej techniki audio XXI wieku.

Kompaktowy wzmacniacz klasy D z wejściem analogowym

Na pierwszy ogień weźmiemy maleńki układ wzmacniacza z różnicowym wejściem analogowym. Układ A21SP16 marki ST Microelectronics, którego schemat blokowy (i aplikacyjny zarazem) pokazano na **rysunku 1**, to prosta konstrukcja pracująca w klasie D, oferująca do 1,4 W mocy wyjściowej przy zasilaniu 5 V i obciążeniu 8 Ω oraz wyposażona w funkcję standby. Całość jest zamknięta w obudowie typu flip-chip o wymiarach zaledwie 1,6×1,6 mm.

Częstotliwość sygnału PWM wynosi 250 kHz, co wprawdzie ułatwia budowę ewentualnego filtra wyjściowego, jednak producent deklaruje, że użycie dodatkowych obwodów dolnoprzepustowych na wyjściu nie jest konieczne. Rzecz jasna, aby uniknąć stosowania dodatkowego układu LC trzeba spełnić kilka warunków. Oprócz uwzględnienia dobrych praktyk projektowania PCB (ekranowanie problematycznych ścieżek płaszczyznami masy, zastosowanie dławików ferrytowych itp.), największe znaczenie dla redukcji zakłóceń EMI ma odległość pomiędzy wzmacniaczem a wyprowadzeniami głośnika. Producent zaleca, by przewód łączący przetwornik z wyjściem A21SP16 miał długość nieprzekraczającą 50 mm – przy dłuższym okablowaniu filtr może okazać się konieczny (o ile chcemy, by nasze urządzenie z powodzeniem przeszło badania EMC w zakresie emisji zakłóceń radiowych).

Układ A21SP16 ma konstrukcję wyjściową typu BTL (Bridge-Tied Load), dzięki czemu głośnik nie wymaga stosowania wyjściowych kondensatorów sprzężenia AC. W najprostszej aplikacji wyprowadzenia głośnika faktycznie są zatem podłączane wprost do pinów wzmacniacza. Jeżeli filtr okaże się konieczny z przyczyn opisanych powyżej, to z powodzeniem można zastosować prosty, jednostopniowy filtr trybu wspólnego (o układzie typowym dla wzmacniacza klasy D z wyjściami typu BTL), pokazany na **rysunku 2**.

Proste wzmacniacze klasy D z wejściem I²S

Rodzina układów MAX98357A/MAX98357B dawnej firmy Maxim (przejętej jakiś czas temu przez Analog Devices) obejmuje niezwykle kompaktowe, 3,2-watowe (@ 4 Ω, 5 V) wzmacniacze klasy D w obudowach WLP (o wymiarach 1,345×1,435 mm) oraz TQFN (3×3 mm). Różnica między wersjami A oraz B leży w interfejsie cyfrowym – MAX98357A obsługuje klasyczne łącze I²S, podczas gdy MAX98357B – strumień TDM z wyrównaniem do lewej. Łącznie obydwa układy pozwalają na obsługę do 35 różnych schematów takowania PCM i TDM i mogą pracować z częstotliwościami próbkowania od 8 kHz do 96 kHz. Użytkownik może na drodze sprzętowej ustalić obsługiwany kanał: do wyboru są zwykłe tryby monofoniczne (lewy lub prawy), a także automatycznie obliczana suma kanałów (a ściślej rzecz ujmując: ich średnia arytmetyczna, tj. prawy/2 + lewy/2). Tak jak to bywa w niewielkich układach o rozbudowanej funkcjonalności (a konfigurowanych jedynie za pomocą pojedynczych pinów), producent zastosował tutaj dość proste, ale skuteczne rozwiązanie: trzy komparatory z fabrycznie ustalonymi napięciami odniesienia współpracują ze wspólnym rezystorem pull-down o wartości 100 kΩ (**rysunek 3**). W trybie TDM możliwe jest natomiast multipleksowanie nawet do czterech kanałów, za co odpowiada stosowna konfiguracja wejścia GAIN_SLOT – patrz **rysunek 4**.

Układy MAX98357A/MAX98357B – podobnie jak ma to miejsce w znacznej części obecnej oferty małych wzmacniaczy scalonych klasy D – mogą pracować bez zewnętrznego filtra LC. Aby zaradzić potencjalnym problemom natury EMC, inżynierowie firmy Maxim zastosowali dwie techniki ograniczania pików widma zaburzeń promieniowanych:

- aktywne sterowanie szybkością narastania/opadania zboczy sygnału na wyjściu modulatora PWM,
- modulacja częstotliwości modulatora (metoda widma rozproszonego), czyli wprowadzanie nieznacznych odchyłek w długości kolejnych okresów sygnału wyjściowego (w sposób nie pogarszający jakości odtwarzanego dźwięku). Odchyłka ta została zresztą jawnie określona przez producenta: wynosi ± 20 kHz przy częstotliwości podstawowej równej 300 kHz.

Opisane rozwiązania pozwoliły na spłaszczenie pików składowej podstawowej oraz harmonicznym wyższych rzędów do poziomu praktycznie nieprzekraczającego 20 dBμV/m oraz – ogólnie rzecz ujmując – zawężenie zakresu istotnego pasma emisji zakłóceń. W praktyce uzyskano odstęp przynajmniej 10 dBμV/m od limitów narzuconych przez normę EN55022B i to w dowolnym punkcie pasma podlegającego badaniom EMC – doskonale widać to na **rysunku 5**, prezentującym wynik badań zaburzeń emitowanych przez system złożony ze wzmacniacza oraz głośnika podłączonego do niego przewodem o długości 12 cali (około 30 cm).

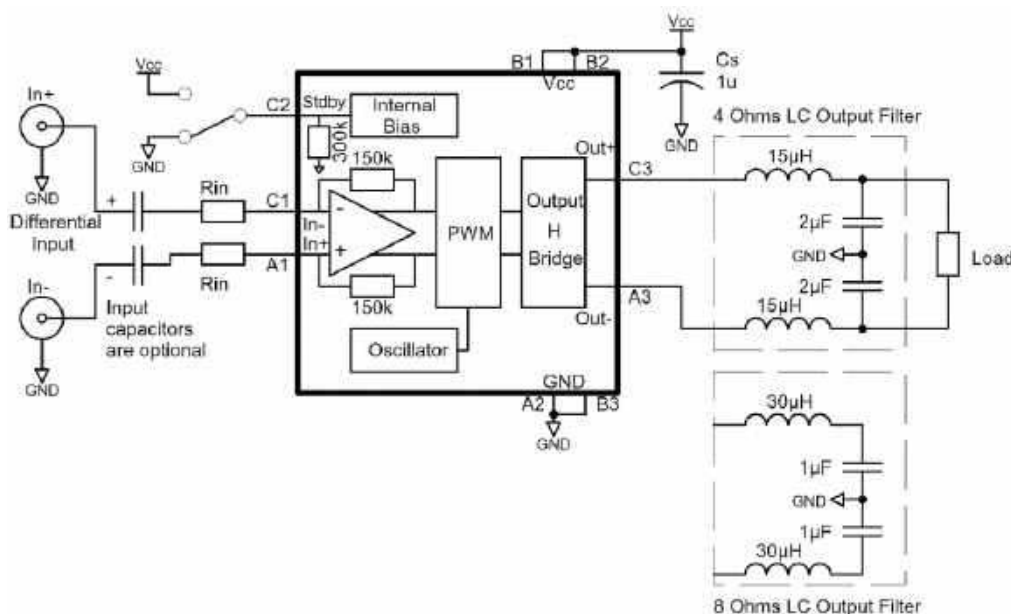
Warto dodać, że zwłaszcza pierwsza z ww. technik może prowadzić do ograniczenia sprawności układu – jak wiadomo, im mniej strome zbocza, tym większe straty przełączania, związane z kluczowaniem

tranzystorów wyjściowych. Finalny efekt jest jednak naprawdę niezły – układ utrzymuje sprawność na poziomie nie gorszym niż 92%.

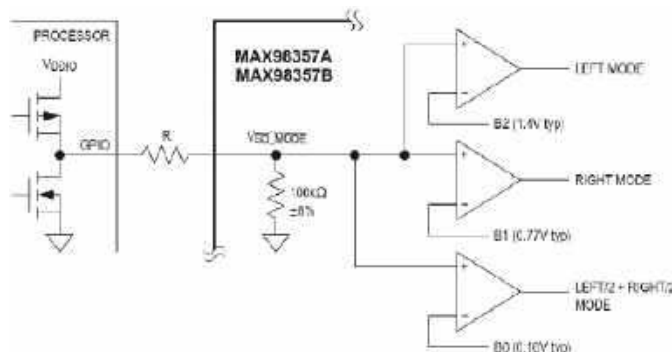
W niektórych urządzeniach (np. słuchawkach bezprzewodowych) dość irytującym zjawiskiem jest dźwięk „kliknięcia” przy włączaniu układu, spowodowany nagłym zasilaniem membrany głośnika. W układach z serii MAX98357A/MAX98357B producent zastosował odpowiednie środki zaradcze, mające na celu eliminację wspomnianego efektu. Na oscylogramie z **rysunku 6** można zobaczyć działanie tego „softstartu” po uruchomieniu układu w wyniku przełączenia linii $\overline{SD_MODE}$ w stan wysoki, czyli podczas wychodzenia z trybu shutdown.

Wzmacniacz wzmacniaczowi nierówny. Układy pracy słuchawek stereo (i nie tylko)

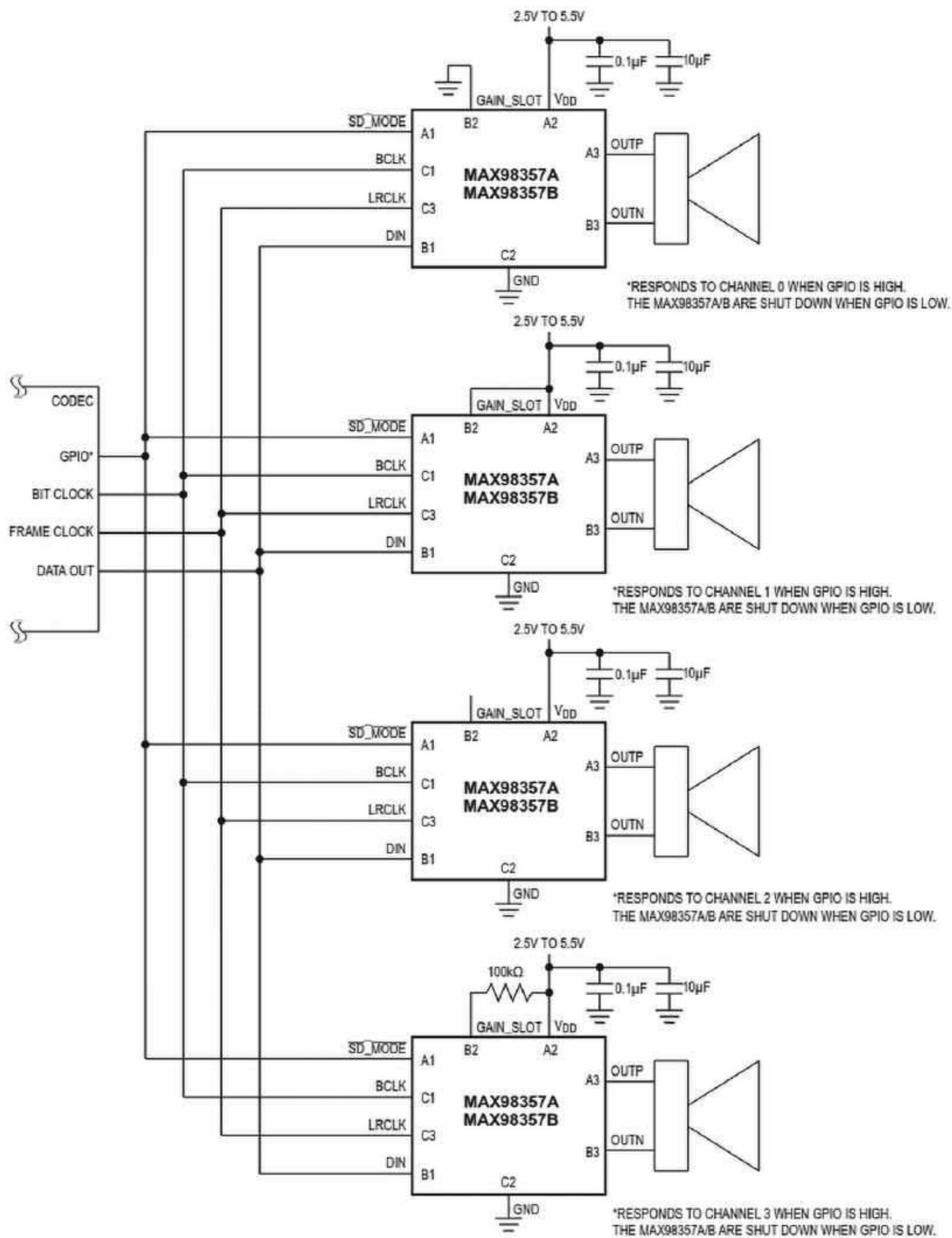
Przerwijmy na chwilę przegląd konkretnych modeli wzmacniaczy scalonych, by przyjrzeć się trzem rozwiązaniom układowym stosowanym do obsługi przewodowych słuchawek stereo. Klasyczny układ, który można spotkać m.in. w rozmaitych wzmacniaczach opartych na scalonych „kostkach” starszego typu, pokazano na **rysunku 7a**. Obydwa kanały są podłączone do gniazda słuchawkowego (najczęściej typu jack 3,5 mm lub jack 6,3 mm) za pośrednictwem kondensatorów szeregowych. Elementy te odgrywają rolę sprzężenia zmiennoprądowego, czyli odcinają składową stałą sygnału wyjściowego. Dzięki temu cały układ wzmacniacza może pracować



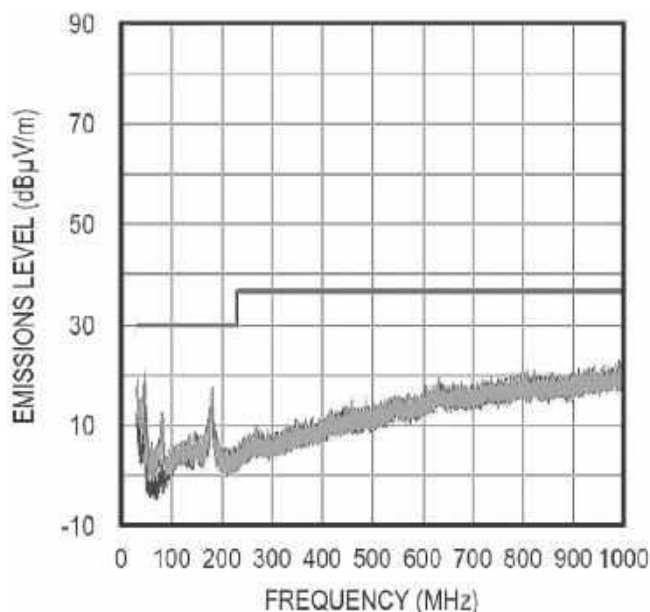
Rysunek 2. Schemat aplikacyjny układu A21SP16 z uwzględnieniem filtrów LC na wyjściach wzmacniacza (filtr jest wymagany przy długości okablowania głośnika przekraczającej 50 mm) [1]



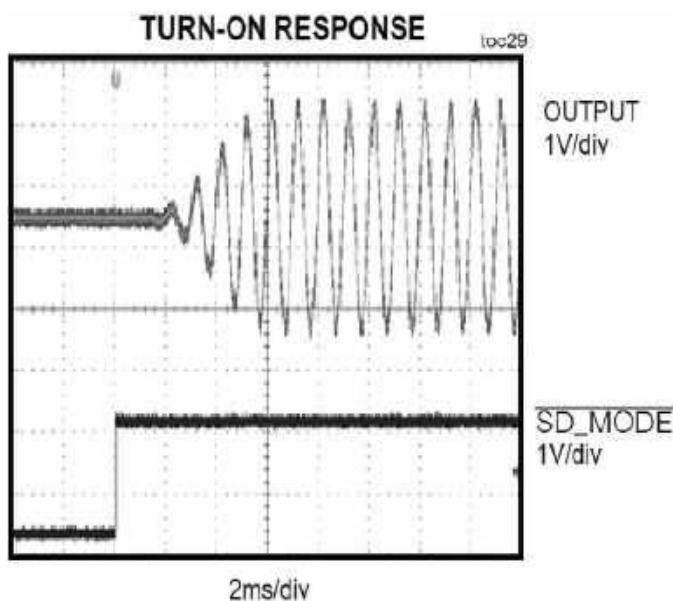
Rysunek 3. Budowa wejścia wyboru kanału w układach MAX98357A/MAX98357B [2]



Rysunek 4. Schemat aplikacyjny układów MAX98357A/MAX98357B w czterokanałowym trybie TDM [2]



Rysunek 5. Spektrum emisyjne wzmacniacza z serii MAX98357A/MAX98357B przy użyciu z głośnikiem podłączonym za pomocą 12-calowego przewodu, bez filtrów LC [2]



Rysunek 6. Ilustracja działania funkcji „softstartu” (Click-and-Pop suppression) [2]

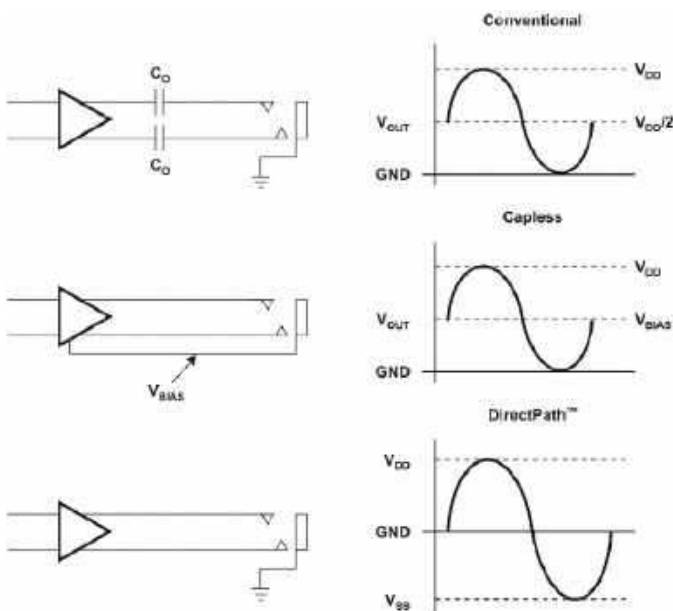
z zasilaniem pojedynczym, a nie z tradycyjnym (choć w zminiaturyzowanych urządzeniach konsumenckich stosowanym coraz rzadziej) zasilaniem symetrycznym. Styk wspólny jest połączony z masą urządzenia, co wydaje się dość oczywiste i naturalne.

Eliminacja kondensatorów sprzęgających pozwala dość drastycznie zmniejszyć wymiary układu – rzecz jasna dotyczy to głównie tych najmniejszych wzmacniaczy scalonych, przy których odpowiednie kondensatory, nawet w wersji SMD, wydają się gigantami. Usunięcie komponentów zdolnych do odcięcia składowej stałej wymaga (w przypadku klasycznych układów wzmacniaczy) innego potraktowania styku wspólnego – konieczne okazuje się bowiem podłączenie go nie do masy, ale do sztucznie wytworzonego potencjału pośredniego (V_{BIAS} na rysunku 7b), równego – a jakże – połowie napięcia zasilającego. Takie rozwiązanie niesie jednak ze sobą szereg niebezpieczeństw i innych wad. Największe ryzyko? Użycie wyjścia słuchawkowego w roli wyjścia liniowego (co ma miejsce w wielu praktycznych sytuacjach) może spowodować, że potencjał „prawdziwej” masy zostanie zwarty z masą sztuczną. Co może pójść nie tak? Chyba każdy odpowie już na to pytanie we własnym zakresie...

Układy pracy pokazane na rysunkach 7a i 7b mają jednak jeszcze jedną, istotną wadę, która okazuje się drastycznym ograniczeniem zwłaszcza w przypadku wzmacniaczy słuchawkowych. Zwróćmy uwagę, że w obydwu przypadkach zakres zmian napięcia waha się od (mniej więcej) potencjału masy do potencjału dodatniej szyny zasilania. W efekcie uzyskujemy 1/4 mocy, którą dałoby się osiągnąć przy zastosowaniu bipolarnego systemu zasilania. Jeżeli wzmacniacz jest zasilany niskim napięciem, a do tego impedancja przetworników jest relatywnie wysoka (16 Ω lub 32 Ω), to otrzymujemy receptę na załóżnie niską moc wyjściową.

W przypadku systemów monofonicznych oraz stereofonicznych (z odrębnymi torami poszczególnych głośników) rozwiązaniem problemu ograniczenia mocy jest zastosowanie końcówki mocy typu mostkowego, czyli tzw. układu BTL (Bridge-Tied Load). Jest to nic innego, jak klasyczny mostek H, dzięki któremu uzyskujemy 2-krotnie wyższe napięcie (i 4-krotnie większą moc) w porównaniu do układu push-pull z obciążeniem dołączonym do masy. Niestety, układu mostkowego nie da się w prosty sposób zastosować w przypadku słuchawek, w których przetworniki mają jedno wyprowadzenie wspólne.

Opisane powyżej problemy doprowadziły do opracowania rozwiązań mających na celu eliminację zarówno kondensatorów sprzęgających z wyjścia wzmacniacza słuchawkowego, jak i konieczności połączenia styku wspólnego obu słuchawek z potencjałem sztucznej masy. W istocie różni producenci stosują różne nazwy handlowe do określenia tego samego triku układowego, polegającego po prostu na wyposażeniu wzmacniacza w niewielką pompę ładunkową. Jej zadaniem jest wytwarzanie ujemnego napięcia zasilającego, dzięki któremu końcówka mocy może pracować dokładnie tak, jak gdyby całość była zasilana napięciem symetrycznym. Rzecz jasna takie rozwiązanie ma sens jedynie w przypadku niewielkiej mocy wyjściowej – dlatego doskonale sprawdza się w układach stereofonicznych słuchawek przewodowych. Układ połączeń wyjść takiego wzmacniacza z gniazdem pokazano na rysunku 7c. Genialnie proste. I równie efektywne.



Rysunek 7. Trzy możliwości podłączenia słuchawek stereofonicznych do wzmacniacza z pojedynczym napięciem zasilającym [3]

Wzmacniacz słuchawkowy ze sterowaniem cyfrowym

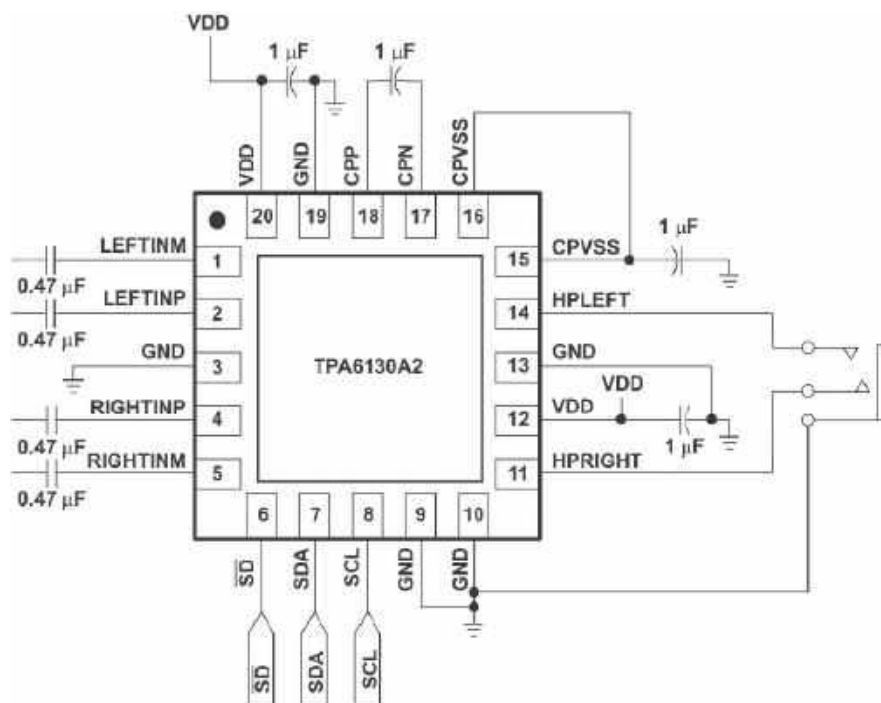
Przykładem wzmacniacza przeznaczanego stricte do aplikacji słuchawkowych może być układ TPA6130A2, oferowany w 4-milimetrowych obudowach QFN oraz 2-milimetrowych DSBGA. Jego schemat aplikacyjny pokazano na **rysunku 8**. Oprócz dwóch kondensatorów odsprężających wymaga on do prawidłowej pracy także dwóch kolejnych pojemności zewnętrznych (odpowiedzialnych za obsługę pompy ładunkowej) oraz wejściowych kondensatorów sprzęgających – łącznie daje to 6 elementów peryferyjnych. I tyle – cała reszta funkcjonalności jest realizowana na poziomie krzemu. Wbudowany układ kontroli wzmocnienia pozwala na ustawienie jednego z 64 poziomów głośności. Dodatkowo, za pomocą tego samego interfejsu I²C, który umożliwia kontrolę mocy wyjściowej, użytkownik jest w stanie sterować wyciszaniem poszczególnych kanałów i trybem pracy, czy też odczytywać aktualny status układu. Dostępna moc wyjściowa wynosi 300 mW przy pracy z 16-omowymi przetwornikami.

Wysoki PSRR układu TPA6130A2 (100 dB min.) pozwala na zasilanie bezpośrednio z akumulatora lub baterii, bez pośrednictwa przetwornicy czy regulatora liniowego. Zasilanie końcówki mocy z niestabilizowanego napięcia zawsze wiąże się jednak z pewnym problemem, polegającym na ograniczeniu mocy wyjściowej w zależności od poziomu rozładowania źródła energii. Z drugiej strony zastosowanie przetwornicy DC/DC wiąże się z niepotrzebnym zużyciem dodatkowej mocy podczas relatywnie cichego odsłuchu – po co bowiem „ładować” we wzmacniacz podwyższone napięcie, skoro doysterowania przetworników na pożądanym poziomie potrzebny jest zaledwie jego ułamek?

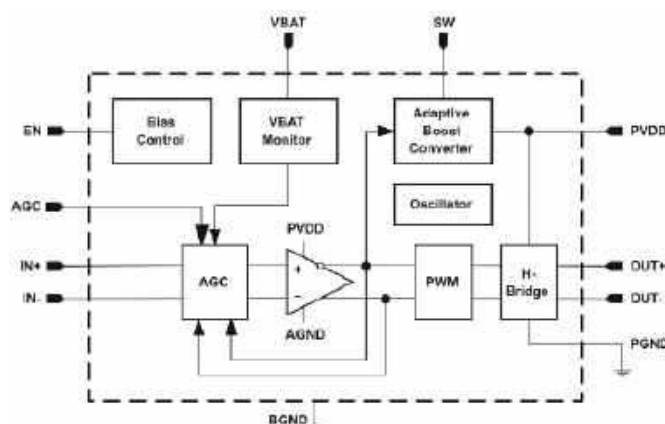
Wzmacniacz z adaptacyjną kontrolą wzmocnienia i napięcia zasilania

Popisem innowacyjności pod względem redukcji zużycia mocy i optymalizacji pracy „na końcówce baterii” może być układ TPA2025D1 (**rysunek 9**). Tym razem mamy do czynienia z układem monofonicznym, w którym końcówka mocy (klasy D w układzie mostkowym) jest zasilana za pomocą adaptacyjnie przestrajanej przetwornicy DC/DC. Na wejściu układu znalazł się blok AGC – co może okazać się zaskakujące, gdyż obwód taki kojarzy się bardziej z torami RF, niż z układami audio. Co jeszcze ciekawsze, jest on przestrajany przez układ monitorujący napięcie zasilania (sic!). To nie chochlik drukarski – projektanci TPA2025D1 wprowadzili takie rozwiązanie po to, by zmniejszyć maksymalny pobór prądu przy krańcowych warunkach rozładowania akumulatora lub zestawu baterii. Gdy napięcie zasilania urządzenia spadnie poniżej pewnego progu (ustalonego za pomocą odpowiedniej konfiguracji pinu AGC), obwody automatycznej kontroli wzmocnienia zaczną stopniowo zmniejszać poziom głośności (nawet o 10 dB przy napięciu V_{BAT} rzędu 2,5 V – patrz **rysunek 10**). Za cenę cichszego odsłuchu możemy więc ochronić co bardziej wrażliwe źródła energii przed zgubnym w skutkach, głębokim rozładowaniem.

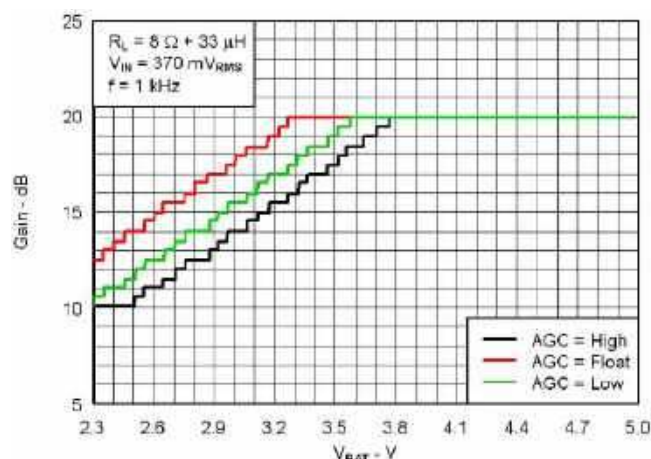
Druga funkcja nosi nazwę *Auto Pass Through* (w skrócie APT) i polega na skokowym załączaniu oraz wyłączaniu przetwornicy w zależności od aktualnego poziomu mocy wyjściowej. Najlepiej



Rysunek 8. Schemat aplikacyjny układu TPA6130A2 [4]

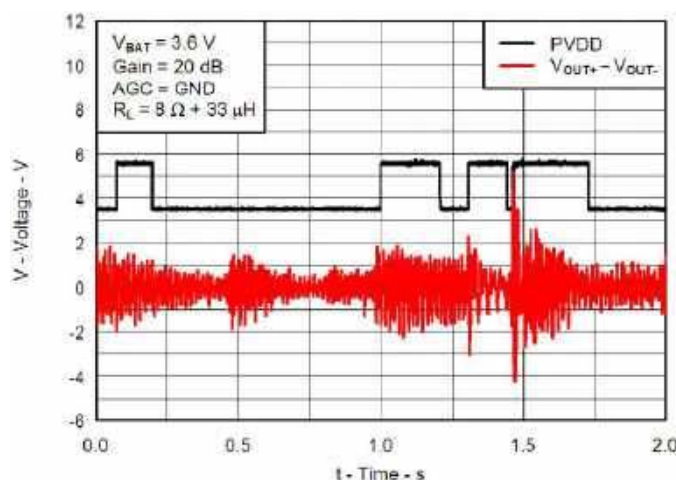


Rysunek 9. Schemat blokowy wzmacniacza TPA2025D1 [5]



Rysunek 10. Ilustracja działania funkcji AGC zaimplementowanej w układzie TPA2025D1 [5]

wyjaśnić to na przykładzie **rysunku 11**. Jak widać, w cichszych partiach odtwarzanego nagrania napięcie zasilania doprowadzone do końcówki mocy jest niższe, a ściślej rzecz biorąc: równe napięciu

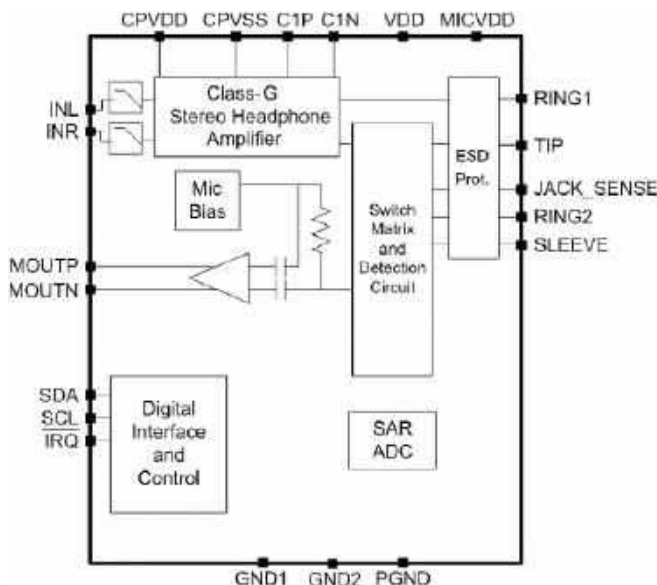


Rysunek 11. Przykład działania adaptacyjnej kontroli napięcia zasilania końcówki mocy układu TPA2025D1 [5]

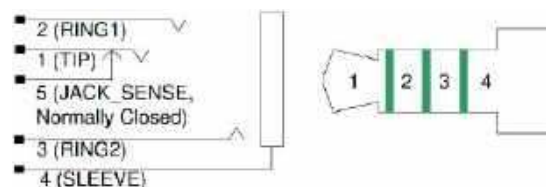
dostarczonemu z zewnętrznego źródła energii. Jeżeli zapotrzebowanie na moc rośnie, kontroler błyskawicznie załącza konwerter DC/DC i podbija napięcie do wartości 5,75 V – dzięki temu układ może odpowiednioysterować wyjście, unikając w ten sposób jego nasycenia (obcinania głośniejszych partii) lub nadmiernych zniekształceń. Jednocześnie – co także ważne – optymalizowany jest pobór mocy, gdyż przetwornica impulsowa pracuje wyłącznie wtedy, gdy naprawdę jest to potrzebne. Maksymalne osiągi układu to moc wyjściowa 1,9 W, dostarczana do głośnika o impedancji 8 Ω i to – uważa! – przy napięciu zasilania całości równym zaledwie 3,6 V.

Front-end do zestawów słuchawkowych

Bardzo ciekawe rozwiązania konstrukcyjne wprowadzili projektanci układu TPA6166A2 (rysunek 12). W odróżnieniu od opisanych wcześniej układów jest to kompletny front-end przeznaczony do obsługi zestawów słuchawkowych (słuchawki + mikrofon) z czterosekcyjnym wtykiem jack 3,5 mm (rysunek 13). Wyprowadzenia układu są dostosowane do typowych gniazd z czterema stykami oraz wbudowanym przełącznikiem umożliwiającym łatwą detekcję faktu podpięcia zestawu do urządzenia. Układ TPA6166A2 umożliwia nie tylko wykrycie podłączenia, ale także maksymalnie upraszcza proces



Rysunek 12. Schemat blokowy wzmacniacza TPA6166A2 [6]

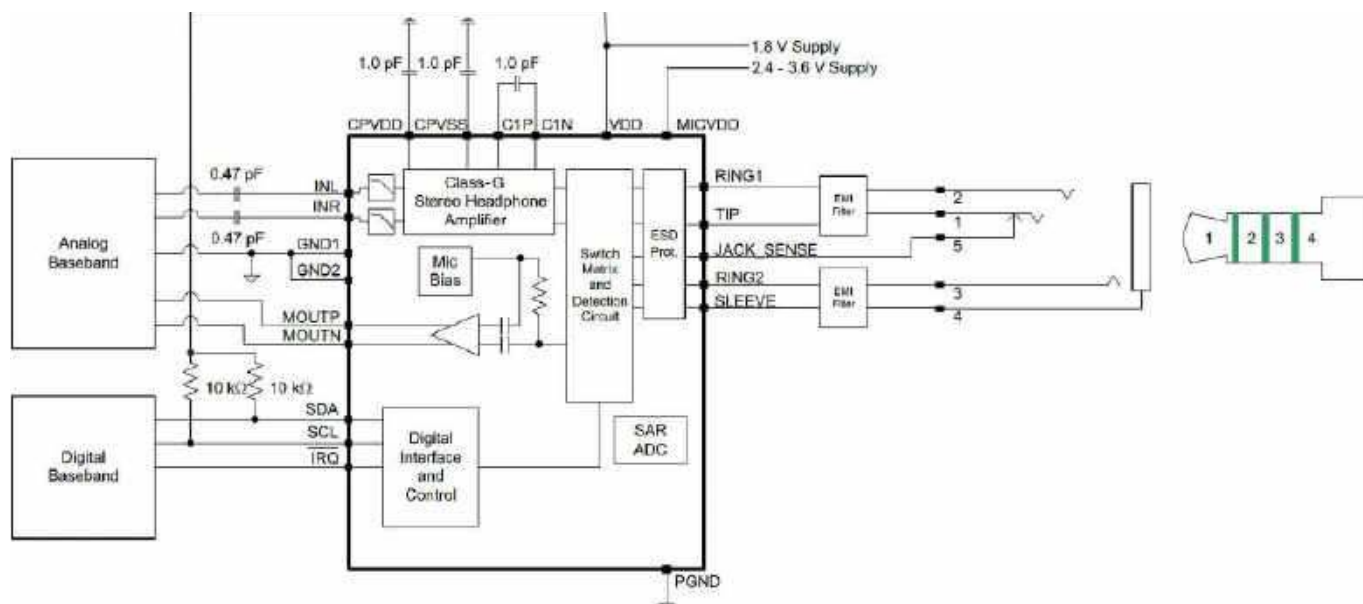


Rysunek 13. Układ wyprowadzeń wtyku jack 3,5 mm typowego zestawu słuchawkowego z mikrofonem [6]

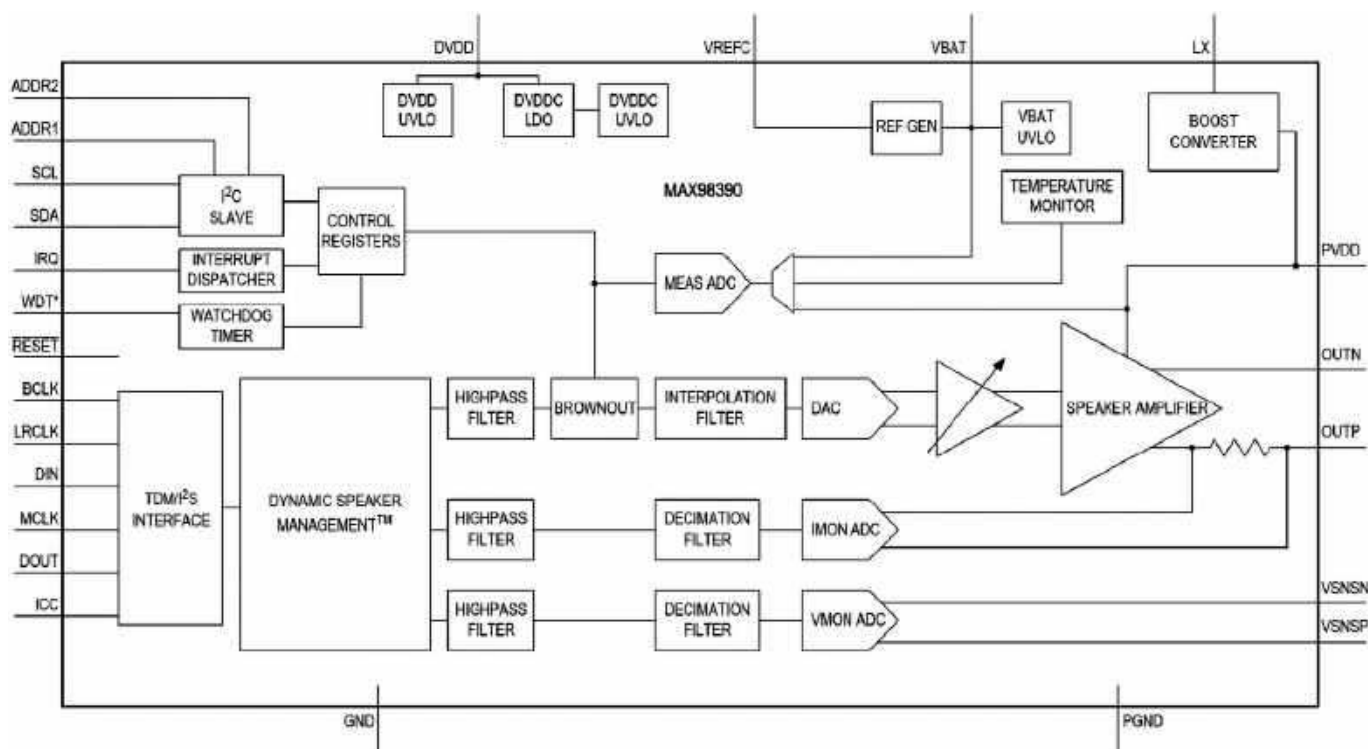
rozpoznawania rodzaju wpiętego akcesorium. Dzięki opracowanym przez inżynierów TI rozwiązaniom w zakresie pomiaru rezystancji pomiędzy odpowiednimi wyprowadzeniami gniazda, nadrzędny kontroler może rozpoznać następujące typy wyposażenia:

- zestaw słuchawkowy stereo z mikrofonem,
- zestaw słuchawkowy mono,
- słuchawki stereo,
- kabel wyjścia liniowego.

Aby całość miała funkcjonalny sens, oprócz wzmacniacza słuchawkowego (klasy G), w strukturze układu znalazł się także zintegrowany



Rysunek 14. Schemat aplikacyjny układu TPA6166A2 [6]



*WDT PIN AVAILABLE ON 36-BUMP PACKAGE ONLY

Rysunek 15. Schemat blokowy wzmacniacza MAX98390 [7]

przedwzmacniacz mikrofonowy o wzmocnieniu ustawianym programowo na 12 dB lub 24 dB. Układ obsługuje ponadto funkcję detekcji naciśnięcia przycisku znajdującego się na wyposażeniu zestawu słuchawkowego – w tym celu mierzona jest impedancja pomiędzy wyprowadzeniami RING2 oraz SLEEVE gniazda jack. Co ciekawe, do dyspozycji użytkownika przewidziany jest w tym celu 10-bitowy przetwornik ADC, który umożliwia detekcję naciśnięcia jednego z kilku przycisków, połączonych w układzie „pasywnej klawiatury”.

Osiągi układu TPA6166A2 w zakresie parametrów audio prezentują się następująco:

- moc wyjściowa: 30 mW/kanal @ 32 Ω,
- THD+N: 1%,
- zakres regulacji głośności: -42 dB...+6 dB,
- szum wyjściowy: 2 μV @ -42 dB,
- sterowane programowo napięcie polaryzujące mikrofon: 2,0 V lub 2,6 V,
- PSRR: 91 dB (kanal słuchawkowy)/92 dB (kanal mikrofonu).

Układ jest dostępny w obudowach WSCP o wymiarach 2,5×2,5 mm. Kompletny układ aplikacyjny, uwzględniający niezbędne minimum elementów dyskretnych w otoczeniu samego wzmacniacza, można zobaczyć na **rysunku 14**. Warto zwrócić uwagę, że – w odróżnieniu od wcześniej opisanych (nowszych) konstrukcji – w tym przypadku wymagane są wyjściowe filtry EMI.

„Inteligentny” wzmacniacz z dynamicznym zarządzaniem głośnikiem

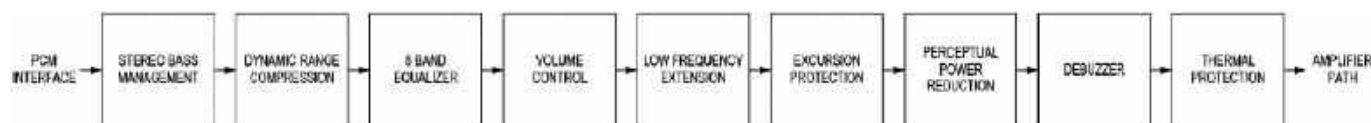
Ten enigmatyczny i nieco zagadkowy podtytuł nie oddaje zbyt dobrze ogromnych możliwości układu MAX98390 (**rysunek 15**). Opisany komponent – dostępny w malutkich obudowach 30-Bump WLP (2,468×2,968 mm) oraz 36-Bump WLP (2,418×2,608 mm) – oferuje zaskakująco wysoką moc wyjściową, dochodzącą do 6,2 W. Wysoki poziom zniekształceń (THD+N do 10% przy maksymalnej mocy) sugeruje, że wzmacniacz – delikatnie rzecz ujmując – niezbyt nadaje się do urządzeń klasy audiofilskiej. Sprawdzi się

za to w rozmaitych zabawkach, urządzeniach IoT czy nawet wbudowanych systemach audio komputerów przenośnych, w których efektywność i moc wyjściowa są zwykle stawiane na pierwszym miejscu, zaś zniekształcenia i tak odchodzą na dalszy plan (zwłaszcza przy zastosowaniu umiarkowanej jakości przetworników).

Już pierwszy rzut oka na rysunek 15 pozwala stwierdzić, że w tym przypadku mamy do czynienia z naprawdę rozbudowaną konstrukcją, w dużej mierze „ucyfrowioną” i kompleksową. Za odbiór sygnałów audio odpowiada hybrydowy interfejs wspierający tryby pracy I2S i TDM, zaś za sterowanie funkcjami (odczyt i zapis rejestrów) – kontroler I2C z dodatkowymi liniami sterującymi oraz sygnałowymi (ustawianie adresu, przerwanie, obsługa watchdoga, reset). Tajemniczy blok dynamicznej kontroli parametrów głośnika znajduje się jeszcze przed DAC, tj. w całości w domenie cyfrowej – na tym etapie bowiem dane odebrane z interfejsu szeregowego są poddawane modyfikacjom mającym na celu:

- podniesienie poziomu ciśnienia akustycznego (wg producenta – nawet do wartości 2,5-krotnie przekraczającej osiągi tego samego głośnika napędzanego przez standardowy wzmacniacz o analogicznych parametrach),
- usunięcie niepożądanych rezonansów, pojawiających się podczas pracy niewielkiego przetwornika przy wysokim poziomie mocy wyjściowej i słyszalnych jako nieprzyjemne brzęczenie,
- rozszerzenie efektywnego pasma akustycznego nawet o dwie oktawy w dół, a innymi słowy: „inteligentne” podbicie poziomu basów,
- redukcję poziomu występowania głośnika przy zachowaniu zbliżonej wierności odtwarzania dźwięku.

Tak intensywna ingerencja w sygnał dźwiękowy ma na celu „wycisnienie” z niewielkiego przetwornika maksimum jego możliwości – nic więc dziwnego, że układu nie należy raczej proponować osobom, dla których jakość odtwarzania ma znaczenie pierwszorzędne. W przypadku przenośnych urządzeń konsumenckich, zabawek czy nawet laptopów nie będzie to jednak problemem, pozwoli za to podnieść jakość dźwięku w stosunku do tradycyjnych



Rysunek 16. Cyfrowy tor sygnałowy odpowiedzialny za realizację funkcji DSM (Dynamic Speaker Management), zaimplementowanej w układzie MAX98390 [7]

implementacji, które ograniczone są nie tylko przez rozmiar i parametry samego głośnika, ale także jego (zwykle nieoptymalne) otoczenie w postaci obudowy urządzenia.

Na rysunku 16 pokazano przebieg cyfrowego toru sygnałowego realizującego funkcję DSM (Dynamic Speaker Management). Dane otrzymane z interfejsu szeregowego są najpierw obrabiane pod kątem basów i poddawane kompresji, a następnie filtrowane za pomocą 8-pasmowego equalizera. Dopiero po nim następuje właściwe dostosowanie poziomu głośności, podbicie basów i szereg ograniczników, łącznie z redukcją częstotliwości rezonansowej głośnika i zabezpieczeniem termicznym końcówki mocy. Dopiero tak przetworzone dane trafiają na końcowy tor sygnałowy z filtrem górnoprzepustowym, filtrem interpolacyjnym i przetwornikiem DAC. Realizacja wymyślnych zabezpieczeń i funkcji dostrajających sygnał do indywidualnych potrzeb danej aplikacji (i to przy zastosowaniu ściśle określonego modelu głośnika) jest możliwa dzięki użyciu dodatkowych dwóch torów wejściowych, odpowiedzialnych za profilowanie mocy (napięcia i prądu) przetwornika. W tym celu producent przewidział osobne torry wyposażone w przetworniki ADC: jeden mierzy prąd na wewnętrznym boczniku, drugi zaś – napięcie (podane z zewnątrz na dedykowane piny układu).

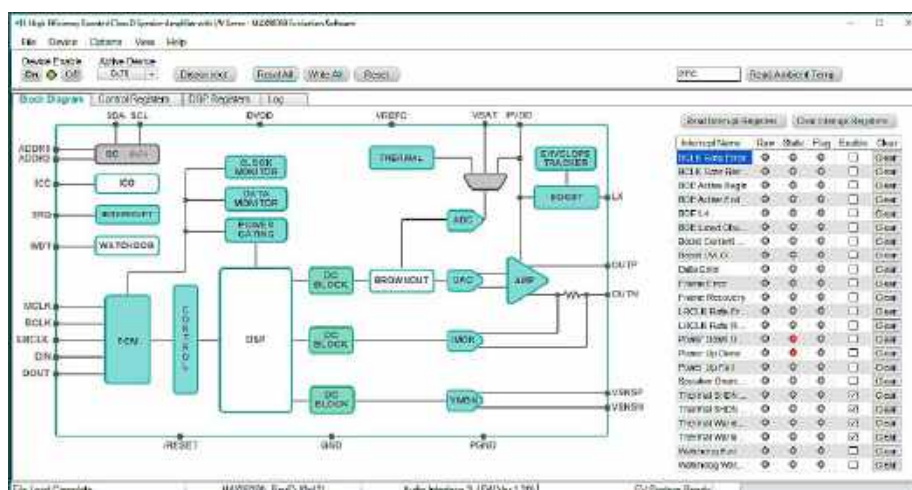
No dobrze... ale na czym polega profilowanie głośnika? Intuicyjnie czujemy, że wymaga ono odpowiedniego ustawienia szeregu parametrów numerycznych (np. tłumienia poszczególnych pasm equalizera). Tylko jak to zrobić? Ręczne ustawianie, jakkolwiek byłoby możliwe, wymagałoby jednak ogromnego nakładu pracy, dostępu do odpowiedniego sprzętu pomiarowego i posiadania niemałego doświadczenia w niuansach audio. Firma Maxim opracowała jednak narzędzia, dzięki którym cały proces można wykonać niemal w 100% automatycznie. Przetestowanie funkcjonalności DSM jest możliwe przy użyciu zestawu ewaluacyjnego MAX98390EVSYS (fotografia 1) oraz specjalnego oprogramowania obsługującego płytkę (rysunek 17). Do profilowania głośnika przygotowano natomiast inną aplikację o nazwie DSM Sound Studio – zdaniem producenta pozwala ona na przeprowadzenie całej procedury w ciągu zaledwie 7 minut.

Zintegrowany wzmacniacz do telefonów komórkowych

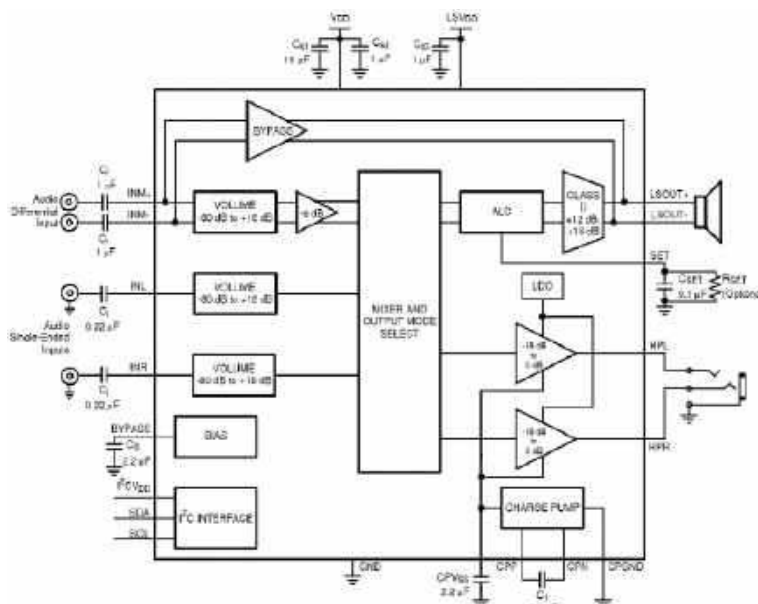
Układ LM49151 powstał przede wszystkim z myślą o smartfonach, ale także komputerach przenośnych, odtwarzaczach MP3 i innych urządzeniach mobilnych. Jego cechą charakterystyczną jest obecność kilku kanałów wejściowych (jednego różnicowego i dwóch monofonicznych pełniących rolę niezbalansowanego wejścia stereo) oraz trzech wyjść (rysunek 18). Pierwsze z nich obsługuje głośnik rozmów (lub inny, niewielki przetwornik o mocy



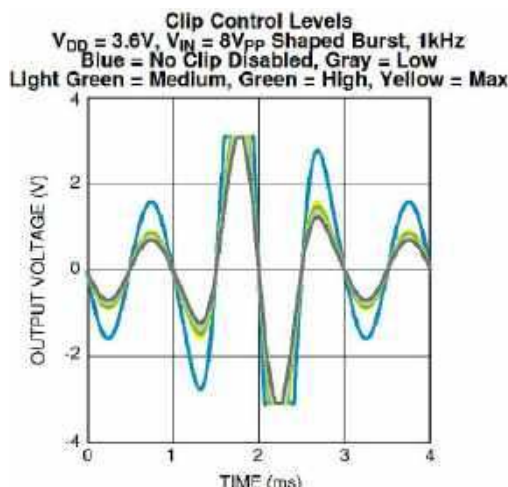
Fotografia 1. Zestaw ewaluacyjny MAX98390EVSYS [8]



Rysunek 17. Oprogramowanie do obsługi płytki MAX98390EVSYS [9]



Rysunek 18. Schemat blokowy wzmacniacza LM49151 [10]



Rysunek 19. Ilustracja działania funkcji ALC wzmacniacza LM49151 [10]

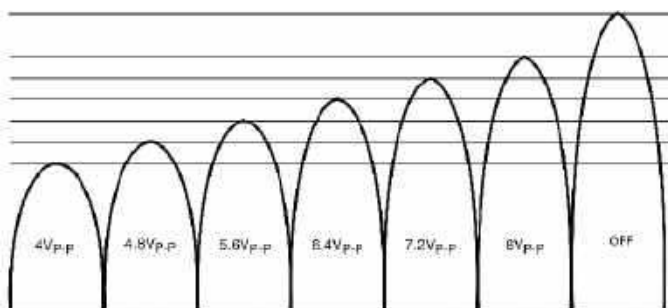
nieprzekraczającej 1,25 W, wbudowany w urządzenie), zaś pozostałe dwa służą do sterowania słuchawkami stereo o mocy do 42 mW/kanal w układzie bez kondensatorów sprzęgających (końcówka mocy jest odniesiona do masy i zasilana z wewnętrznej pompy ładunkowej).

W torze „dużej mocy” obecny jest układ automatycznego przestrajania wzmocnienia (ALC), odpowiedzialny za zabezpieczanie głośnika i ochronę przed obcinaniem wierzchołków sygnału o wyższej amplitudzie (rysunek 19). Producent przewidział 6 różnych poziomów dopuszczalnej amplitudy napięcia wyjściowego oraz siódmy poziom, czyli po prostu tryb wyłączenia funkcji ALC – poszczególne poziomy różnią się od sąsiednich o 0,8 V, co – począwszy od 4 Vpp – daje zakres aż do 8 Vpp. Graficzną ilustrację tego przedziału można zobaczyć na rysunku 20 – wyboru dokonuje użytkownik ustawiając za pośrednictwem interfejsu PC odpowiednie bity jednego z rejestrów konfiguracyjnych układu.

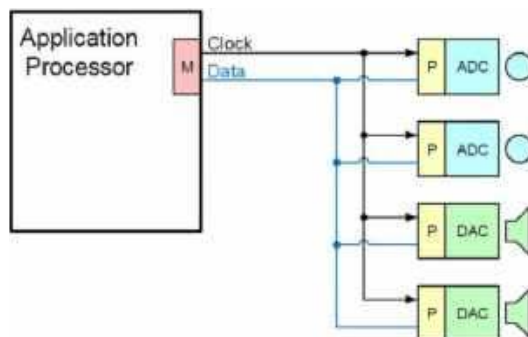
Miniaturowy wzmacniacz z interfejsem SoundWire i zoptymalizowaną obudową WLP

Powróćmy jeszcze na chwilę do miniaturowych (i z pozoru tylko prostych) wzmacniaczy z wejściami cyfrowymi. Opisane przeze mnie na początku artykułu układy oferowały możliwość pracy z sygnałami przesyłanymi za pośrednictwem popularnych interfejsów I²S lub TDM. Nie są to jednak jedyne opcje dostępne dla projektantów cyfrowych systemów audio.

W 2014 roku znane konsorcjum MIPI opublikowało nowy protokół cyfrowego audio, określany mianem SoundWire. Jego celem było upowszechnienie skalowalnego, uniwersalnego interfejsu wymiany danych pomiędzy różnymi układami i urządzeniami, zwłaszcza na pokładzie zaawansowanych systemów wbudowanych: urządzeń mobilnych, komputerów przenośnych itp. Choć w warstwie fizycznej interfejs opiera się na zaledwie dwóch liniach



Rysunek 20. Predefiniowane poziomy odcięcia funkcji ALC wzmacniacza LM49151 [10]

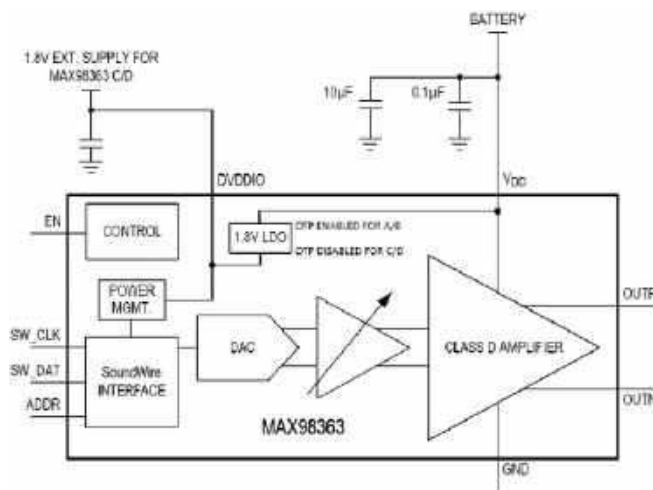


Rysunek 21. Schemat blokowy magistrali MIPI SoundWire z procesorem aplikacyjnym i czterema układami podrzędnymi: dwoma przetwornikami ADC oraz dwoma DAC [11]

(zegara i danych), to w żadnym wypadku nie należy go utożsamiać z I²C – w przeciwieństwie do niego (a także I²S) jest to bowiem łączy na wskroś uniwersalne. Umożliwia zarówno wymianę właściwych danych dźwiękowych, jak i wszelkiego rodzaju informacji sterujących, komend czy identyfikatorów. Co więcej, informacje są przesyłane w trybie DDR, czyli na obydwu zboczach sygnału zegarowego, zaś możliwość współpracy wielu układów na tej samej magistrali (rysunek 21) sprawia, że MIPI SoundWire pozwala na budowę złożonych systemów za pomocą niezwykle prostego layoutu PCB.

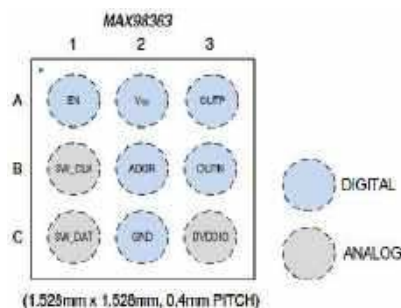
Układ MAX98363 (rysunek 22) jest doskonałym przykładem banalnie prostego w implementacji wzmacniacza z interfejsem SoundWire. W obudowie WLP o boku zaledwie 1,528 mm znalazła się nie tylko końcówka mocy pracująca w klasie D z mocą wyjściową do 3,2 W (@ 4 Ω, 5 V, THD+N 10%), ale także programowalny blok slew-rate (minimalizujący zakłócenia elektromagnetyczne emitowane przez impulsowe wyjście wzmacniacza), układy redukcji artefaktów dźwiękowych (*click-and-pop reduction*), a nawet... wbudowany generator tonów przeznaczony do przeprowadzania fabrycznych testów urządzeń na linii produkcyjnej (!). W strukturze malucha uwzględniono także stabilizator LDO o napięciu 1,8 V oraz analogowy układ regulacji wzmocnienia.

Konstruktorom układu MAX98363 należą się słowa uznania za bardzo przemyślany układ pinów obudowy (rysunek 23). Wszystkie wyprowadzenia oprócz linii ustalającej adres układu znajdują się na obrzeżach obudowy, zatem wyprowadzenie z nich ścieżek (fanout) jest banalnie proste. A co ze środkowym pinem ADDR? Ten należy podłączyć do jednego z dwóch potencjałów zasilania (VDD lub GND), bezpośredniego lub poprzez rezystor 100 kΩ. Pozostaje też piąta opcja, czyli pozostawienie



Rysunek 22. Schemat blokowy wzmacniacza MAX98363 z interfejsem SoundWire [12]

wyprowadzenia wi-
szącego w powietrzu.
Co to oznacza? W prak-
tyce trzy z pięciu dostęp-
nych adresów (ID układu)
można ustawić bez
konieczności wyprowa-
dzania ścieżki od pinu
ADDR poza obrys obu-
dowy. Dzięki temu moż-
liwe jest użycie układu
także w tych aplikacjach,
które ze względów ekono-
micznych nie mogą być zrealizowane z użyciem niezwykle kosztow-
nych przelotek typu via-in-pad. Taki ukłon w stronę konstruktorów PCB
z pewnością docenią wszyscy ci spośród naszych Czytelników, któ-
rzy choć raz przeklinali nieprzemysłany pinout układu w obudowie
WLP czy BGA, zmuszający ich do stosowania obwodów drukowanych
o znacznie droższej technologii, niż faktycznie byłoby to konieczne,
gdyby tylko producent układu głębiej zastanowił się nad optymaliza-
cją rozkładu wyprowadzeń.



Rysunek 23. Układ wyprowadzeń układu MAX98363 [12]

Wzmacniacz z wbudowanymi obwodami ANC

Na początku artykułu wspominałem już o technologii aktywnej
redukcji szumu (ANC), która zrewolucjonizowała elektronikę kon-
sumencką i na stałe weszła do naszego codziennego życia. W naj-
wyższej klasy słuchawkach bezprzewodowych funkcja ANC jest
realizowana przez zaawansowane procesory DSP, „karmione” da-
nymi z zestawu wbudowanych mikrofonów. Prostsza implementa-
cja może być jednak wykonana w oparciu o odpowiednio dostrojony
układ analogowy – dokładnie takie podejście zastosowali projektanci
układu AS3435 marki ams OSRAM (rysunek 24). Niestety, producent
nie upublicznia szczegółowej noty katalogowej, a opis znajdujący się
na stronie internetowej jest krótki i dość mało konkretny. Wiadomo
jednak, że układ jest dostosowany do pracy z bardzo niskim napię-
ciem zasilania (w zakresie 1,0 V...1,8 V) oraz że oferuje moc wyjściową
do 120 mW. Według rozbieżnych informacji oferuje 20-, a nawet 30-de-
cybelowe tłumienie hałasu, głównie w zakresie niższych częstotli-
wości. Jako ciekawostkę dodam, że jeden z blogerów zidentyfikował

omawiany wzmacniacz na PCB słuchawek bezprzewodowych Mid
A.N.C firmy Marshall. Krótki research zasobów internetowych poka-
zał, że – choć funkcjonalności ANC, oferowanej przez producenta kul-
towych pieców gitarowych, daleko do tej znanej np. ze słuchawek
Sony – układ AS3435 radzi sobie i tak całkiem nieźle, pomimo braku
skomplikowanych, autorskich algorytmów DSP czy zaawansowa-
nych procesorów sygnałowych.

Podsumowanie

Jak widać, rynek małych wzmacniaczy audio nie jest ani są-
py, ani mało interesujący. Wręcz przeciwnie – producenci półtora-
i dwumilimetrowych kostek mają do zaoferowania naprawdę wiele,
a głównymi trendami, które wybijają się na czoło wyścigu, są:

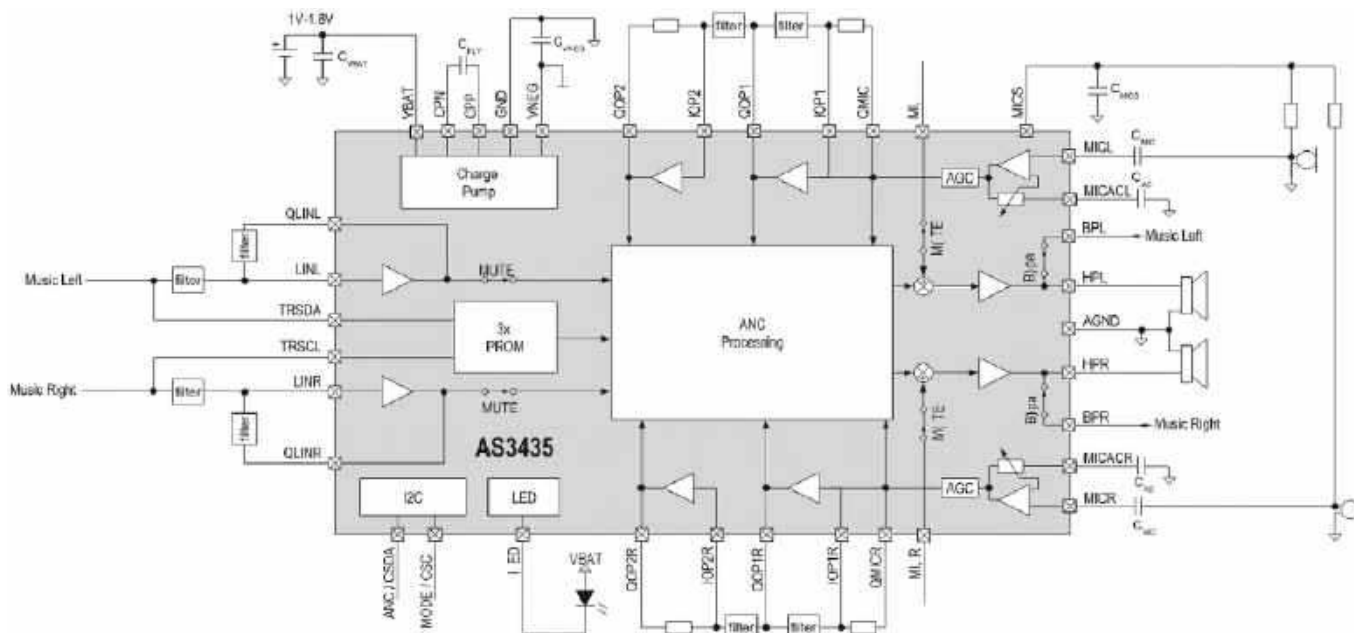
- dążenie do możliwie jak najprostszej implementacji układowej z absolutnie minimalną liczbą komponentów dyskretnych,
- stosowanie technik redukujących zakłócenia EMI i niwelujących konieczność stosowania filtrów wyjściowych we wzmacniaczach klasy D,
- wyposażanie wzmacniaczy w liczne funkcje mające na celu aktywną poprawę jakości dźwięku i/lub optymalizację poboru mocy.

Co przyniesie przyszłość? Z pewnością możemy spodziewać się po-
wolnego hamowania pędu do miniaturyzacji, gdyż zwyczajnie... je-
steśmy raczej blisko granicy. Układy o rozmiarach rzędu 1,5 mm są już
tak małe, że z powodzeniem nadają się nawet do słuchawek TWS.
Być może scalone wzmacniacze słuchawkowe przyszłości będą co-
raz częściej implementowały funkcjonalność brzegowej AI, co po-
zwoli w sposób prosty, a zarazem bezpieczny, optymalizować jakość
dźwięku odtwarzanego z małych przetworników? Czas pokaże.
A my na pewno będziemy trzymać rękę na pulsie.

Jakub Nowicki, EP

Źródła:

- [1] <https://t.ly/6g46v>
- [2] <https://t.ly/1DCEJ>
- [3] <https://t.ly/BqI3f>
- [4] <https://t.ly/IOYh4>
- [5] <https://t.ly/o3A8d>
- [6] <https://t.ly/VtjdL>
- [7] <https://t.ly/SUMaA>
- [8] <https://t.ly/xbKLw>
- [9] <https://t.ly/zkU-u>
- [10] <https://t.ly/uXI57>
- [11] <https://t.ly/BYsRV>
- [12] <https://t.ly/rzmpn>
- [13] <https://t.ly/8hyYM>





**Mnóstwo świetnych
materiałów o audio znajdziesz na
<https://ep.com.pl>**

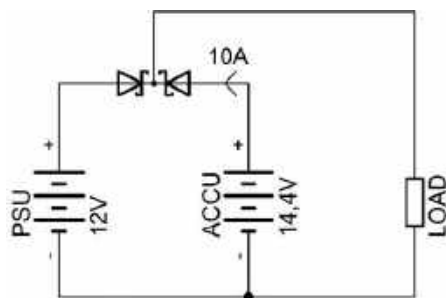


Ograniczenia diody idealnej

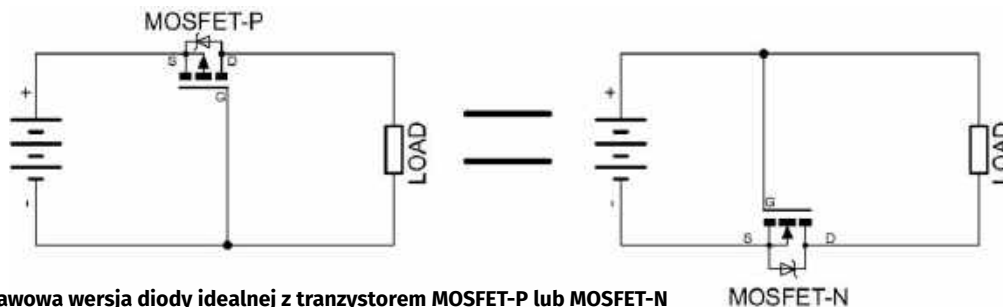
Typowe diody półprzewodnikowe mają spadek napięcia w kierunku przewodzenia sięgający nawet jednego wolta. Można obejść ten problem stosując odpowiednio podłączone tranzystory MOSFET, jednak takie układy wcale nie stanowią doskonałego remedium na każdą przypadłość. O jakich sytuacjach mowa?

Jako przykład do dalszej analizy proponuję następującą sytuację: dwa źródła napięcia o wartości około 12 V (akumulator żelowy, doładowywany przez OZE, a także zasilacz sieciowy) mają za zadanie zasilac system kamer, pobierający prąd o natężeniu około 10 A. Źródła napięcia są dwa, odbiornik zaś jeden – trzeba wybrać to, które aktualnie ma najwyższe napięcie i to z niego czerpać energię. Akumulator 12 V ma z reguły napięcie wyższe niż zasilacz 12 V, bo wynoszące około 14,4 V, więc przez większość czasu to ten pierwszy będzie dostarczał energii. Dopiero kiedy ulegnie znacznemu rozładowaniu, układ przełączy system kamer na pracę z zasilacza.

Najprostszym rozwiązaniem byłoby użycie dwóch diod, włączonych szeregowo z każdym ze źródeł, jak na **rysunku 1**. Przewodzić będzie ta,



Rysunek 1. Diodowy selektor źródeł zasilania

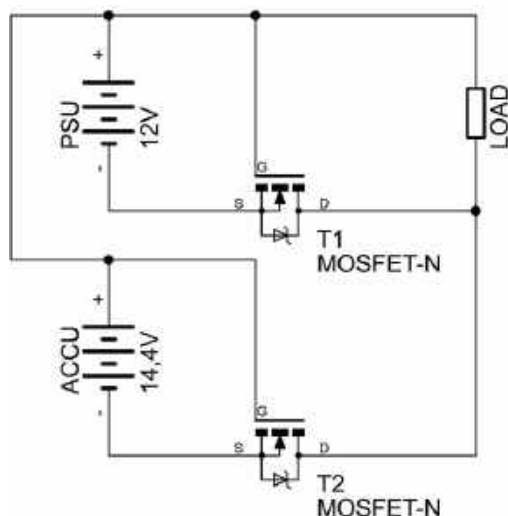


Rysunek 2. Podstawowa wersja diody idealnej z tranzystorem MOSFET-P lub MOSFET-N

która ma najwyższy potencjał anody, druga zaś ulegnie zatkaniu. Taki selektor jest prosty, tani, szybki, trwały, zrozumiały i... energochłonny. Bowiem diody Schottky'ego, nawet te o niskim dopuszczalnym napięciu wstecznym, cechują się napięciem przewodzenia na tyle wysokim, że nie sposób je w takiej aplikacji pominąć. Na przykład element typu VS-40CPQ035-N3 firmy Vishay, który zawiera dwie diody połączone katodami, może przewodzić prąd o natężeniu do 20 A na jedno wyprowadzenie (czyli 40 A przez obie struktury naraz) przy napięciu wstecznym do 35 V. Spadek napięcia na przewodzącej diodzie może wynosić nawet 0,4 V, przez co będzie się na tym elemencie bez przerwy traciła moc rzędu 4 W. Trzeba też uwzględnić fakt, iż na centralkę sterującą pracą kamer trafi napięcie niższe o co najmniej 0,4 V.

Niby niewiele, ale czemu mamy godzić się na takie straty? Poszukajmy czegoś innego, co zmniejszy ilość marnowanej mocy. W literaturze (albo raczej – na różnych internetowych „mądrych” stronach) można znaleźć układ tak zwanej diody idealnej, która rzekomo doskonale nadaje się do zabezpieczenia wejść układów przed odwrotnym podłączeniem do nich napięcia zasilającego – **rysunek 2**. Jeżeli polaryzacja jest prawidłowa, wówczas tranzystor MOSFET-P otwiera się, ponieważ potencjał bramki jest niższy niż źródła. Przy odwrotnym podłączeniu nie nastąpi utworzenie kanału w strukturze tranzystora, bowiem bramka będzie miała potencjał dodatni względem źródła. Można zrobić jeszcze lepszy numer i w przewód ujemny włączyć tranzystor MOSFET-N, który będzie działał analogicznie, ale zaoferuje niższą rezystancję otwartego kanału z uwagi na większą ruchliwość nośników w kanale. W praktycznej realizacji warto zadbać o ograniczenie dopuszczalnego napięcia bramka-źródło, jednak na tym etapie analizy można ten aspekt pominąć.

Ktoś mógłby pomyśleć – mamy jedną diodę, to weźmy dwie i gotowe! Na przykład, taki układ jak z **rysunku 3** jest prosty i niedrogi, gdyż zawiera jedynie dwa tranzystory typu IRL7833. Ich rezystancja otwartego kanału nie przekracza 3,8 mΩ, więc w stanie przewodzenia moc

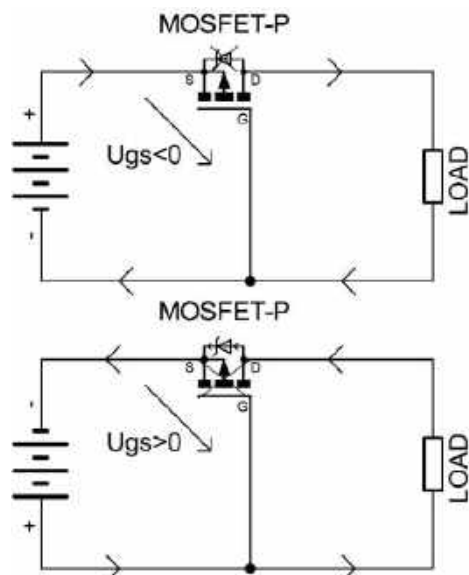


Rysunek 3. Selektor zasilania z dwóch diod idealnych

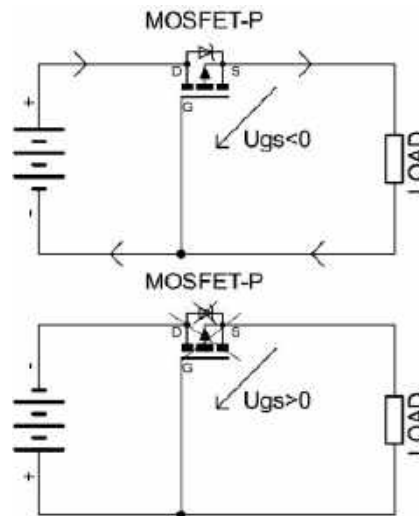
tracona wyniesie niecałe 0,4 W. Bajeczne osiągi, nawet radiator nie będzie potrzebny! Tylko czy aby na pewno to dobre rozwiązanie? Taki pomysł został dawno temu wdrożony w życie (choć na innych tranzystorach), zaś najbardziej wart zapamiętania był fakt, że... nie działał.

Podstawowa przyczyna to brak sterowania tranzystorami. Zasilacz polaryzuje „swoją” T1 napięciem 12 V, przez co jego kanał przewodzi. Ale akumulator również nie próżnuje i aplikuje „swojemu” T2 napięcie rzędu 14,4 V między bramkę, przez co on również przewodzi. Z lekcji podstaw elementów półprzewodnikowych wiemy z kolei, że nośniki w raz utworzonym kanale mogą płynąć w obie strony, ponieważ wszystkie trzy obszary (drenu, źródła i kanał między nimi) mają ten sam typ przewodnictwa, więc na żadnej z dwóch granic (źródło-kanał lub dren-źródło) nie tworzy się złącze. Mamy zatem zasilacz zwarty z akumulatorem przy użyciu dwóch elementów półprzewodnikowych, które nic sobie nie robią z faktu, że mają być tutaj jakimiś diodami. Dlatego też akumulator „pcha” prąd w zasilacz (co może go uszkodzić), jak również zasilacz doładowuje akumulator (co niepotrzebnie zużywa energię z sieci oraz w sposób niekontrolowany ładuje akumulator).

Tutaj nie da się zaradzić nic mądrego, bo ten układ jest po prostu do kitu. Do mądrzejszego sterowania tymi tranzystorami trzeba zastosować bardziej rozbudowany układ, jak chociażby AVT5658. Warto jednak zwrócić uwagę na jeszcze jedną niedoskonałość układu znanego jako dioda idealna. Przeanalizujemy jeszcze raz układ z rysunku 2, dla uproszczenia tylko z tranzystorem MOSFET-P, w dwóch sytuacjach: ze źródłem zasilania wpiętym prawidłowo oraz odwróconym



Rysunek 4. Błędne działanie układu diody idealnej



Rysunek 5. Skorygowana wersja układu diody idealnej

– **rysunek 4.** W sytuacji „prawidłowej” przewodzi tranzystor, bo jego napięcie bramka-źródło jest ujemne. Z kolei w sytuacji „nieprawidłowej”, prąd płynie przez diodę tworzącą się w strukturze tranzystora i głęboko w nosie ma jakiś zatkany tranzystor – ten układ po prostu nie działa, ponieważ ktoś radośnie zlekceważył diodę między drenem a źródłem (dokładniej: między drenem a podłożem zwartym ze źródłem) w tranzystorze MOSFET.

Cóż więc robić, jak to ratować? Trzeba odwrócić tranzystor tak, aby problematyczna dioda stała się naszym sprzymierzeńcem. Tak też jest na **rysunku 5**, który na pierwszy rzut oka działa nieco dziwnie, bo tranzystor polaryzowany jest napięciem nie z akumulatora, ale odkładającym się na obciążeniu. Tego z kolei dostarcza w pierwszym kroku owa niesforna dioda. Kiedy napięcie na obciążeniu wzrośnie powyżej napięcia progowego tranzystora (bo jest ono równe napięciu bramka-źródło), zacznie się on otwierać, co przyspieszy dalszy wzrost napięcia na obciążeniu i tak aż do całkowitego otwarcia – mamy więc mechanizm dodatniego sprzężenia zwrotnego, który jest tutaj pomocny, bowiem przyspiesza załączanie. Przy odwrotnym podłączeniu zasilania nie ma tego problemu, gdyż dioda jest spolaryzowana zaporowo, zaś napięcie bramka-źródło: dodatnie.

Niestety, nawet zaadaptowanie tej skorygowanej wersji diody idealnej do schematu z rysunku 3 nie spowoduje, że zacznie on działać prawidłowo, gdyż dalej oba tranzystory będą przewodziły, sterowane jednocześnie przez napięcie odkładające się na obciążeniu.

W większości innych sytuacji układowych taka idealna dioda będzie wystarczająca, ale nie zawsze. Wadą jest jej czas reakcji w przypadku przełączenia ze stanu przewodzenia do zatkania. Trzeba zapamiętać, że tranzystorem steruje napięcie na obciążeniu, a nie zasilające. Może zatem wystąpić taka sytuacja, w której:

- układ jest zasilany napięciem o prawidłowej polaryzacji, kondensatory w obciążeniu są naładowane,
- biegunowość źródła zasilającego zmienia się gwałtownie na nieprawidłową,
- prąd dalej płynie (w przeciwną stronę!), ponieważ kondensatory obciążenia dalej utrzymują tranzystor w stanie otwarcia.

Z tego powodu układ diody idealnej (w swoim najprostszym, tutaj pokazanym) wydaniu nie znajduje zastosowania w prostownikach. Szkoda, bo oferuje spadek napięcia znacznie niższy w porównaniu do diod półprzewodnikowych. Można jednak pokusić się o zastosowanie specjalizowanego układu scalonego do obsługi tranzystorowego mostka Graetza, co rozwiązuje problem aktywnego sterowania tranzystorami.

Michał Kurzela, EP

Bibliografia:

<https://www.ti.com/lit/an/slvae57b/slvae57b.pdf>

Druk 3D w służbie elektroniki (5)

W poprzedniej części przygotowaliśmy prostą obudowę do gotowego projektu, przy okazji poznając dość dokładnie program CAD o nazwie Autodesk Fusion. W tej części omówimy kwestię składania obudowy i montażu elementów w jej wnętrzu, głównie za pomocą wkrętów metrycznych, nakrętek i wtapianych wkładek gwintowanych oraz zapoznamy się z podstawami projektowania estetycznych i trwałych obudów. Opierając się na praktycznym przykładzie obudowy wykonanej docelowo z tworzywa PETG.

Obudowa, którą będziemy projektować, przeznaczona jest do zasilacza precyzyjnego. Ze względu na jej wielkość i konieczność solidnego trzymania kilku cięższych elementów, jak zasilacz impulsowy czy radiator, obudowa będzie wykonana z indywidualnie drukowanych paneli, które później zostaną złożone razem. Część z nich będzie związana ze sobą na stałe, ale tam, gdzie może być konieczny demontaż (zasilacz jest na etapie prototypowania), zastosowane zostaną wkładki gwintowane i śruby M2,5. Dzięki temu obudowa może być wielokrotnie otwierana i zamykana bez ryzyka uszkodzenia otworów montażowych. Na początek jednak zapoznajmy się z innymi metodami.

Nakrętki wciskane, wsuwane i ukryte oraz inne opcje

Ważną zaletą masowej produkcji jest standaryzacja wymiarów bardzo wielu komponentów. Elektronik zdaje sobie z tego sprawę na przykładzie standardowych obudów komponentów, ale standaryzacja obejmuje też wiele innych części, od wkrętów, podkładek i nakrętek, przez łożyska, silniki elektryczne, koła zębate i przekładnie, aż po wtyki, gniazda, a nawet magnesy. Dzięki temu projektant nie musi się zastanawiać, skąd weźmie pasujący komponent, a producent może mieć listę kilku czy kilkunastu dostawców potrzebnych części. W naszym konkretnym przypadku skupimy się na typowych gwintach metrycznych, czyli standardzie ISO obowiązującym od 1947 roku. Elementy te oznaczane są literą M, po której podana jest średnica gwintu i są to najbardziej powszechne na świecie elementy łączące ogólnego przeznaczenia. **Tabela 1** zestawia istotne dla nas wymiary śrub i nakrętek w rozmiarach najczęściej spotykanych w druku 3D.

W przypadku nakrętek i łbów sześciokątnych podano średnicę, na której sześciokąt jest opisany, maksymalną szerokość łba lub



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

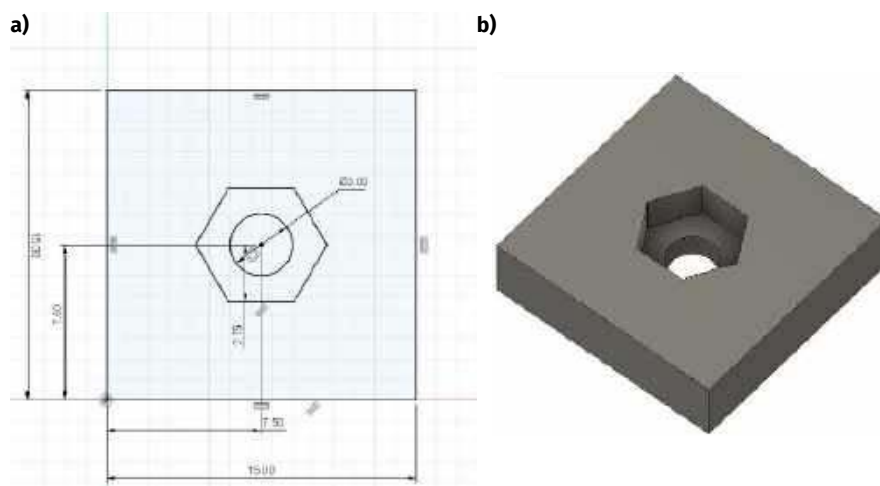
nakrętki oraz grubość nakrętki. Ten ostatni wymiar przyda się, gdy będziemy chcieli wsunąć nakrętkę w otwór, z którego sama nie wypadnie.

Najprostsze rozwiązanie to zaprojektowanie wgłębienia na nakrętkę. Na **rysunku 1a** pokazano prosty szkic. Kwadratowy element o wymiarach 15×15 mm zawiera otwór 3 mm pod wkręt M3. Za pomocą narzędzia „Wielokąt opisany” dodano również wyśrodkowany sześciokąt, którego promień został wpisany jako 5,5/2 mm. By jego orientacja była poprawna, jeden z boków został ograniczony wiązaniem „Równoległy” względem jednego z boków kwadratu. W kolejnych krokach wyciągnięto zewnętrzną część kształtu na 4 mm, a potem wewnętrzną na 4–2,4 mm=1,6 mm, czyli na tyle, by nakrętka licowała z górną z powierzchnią. **Rysunek 1b** pokazuje wynikową bryłę, a **fotografia 1** – gotowy element z nakrętką. Rozwiązanie to jest łatwe w implementacji, a wytrzymałość na wyrwanie nakrętki zależy od liczby warstw górnych i dolnych, gdyż to one przenoszą naprężenia. Grubość ścianek też ma znaczenie dla wytrzymałości całego elementu, a typ i gęstość wypełnienia określają sztywność i zdolność przenoszenia naprężeń między warstwami. Wadą tego rozwiązania jest fakt, iż nic (poza tarcie o ścianki boczne wgłębienia) nie trzyma nakrętki na miejscu, przez co nakrętka będzie wypadać, jeśli element zostanie wydrukowany ze zbyt dużą tolerancją. Z kolei zbyt mały luz montażowy może utrudnić wciśnięcie nakrętki na miejsce.

Następne rozwiązanie sprawdza się najlepiej w bardziej litych elementach, gdyż wymaga większej grubości elementu. **Rysunek 2a** prezentuje zmodyfikowany szkic z rysunku 1a. Dodano linię konstrukcyjną od środka otworu na wkręt do bocznej krawędzi elementu. Za pomocą narzędzia „Odsunięcie” dodano do tej linii dwie równoległe linie w odstępach 0,5 mm. Stworzą one kanał, który przyda się w razie konieczności wymiany nakrętki. Za pomocą narzędzia „Rozwijaj” wydłużono dwie linie boków nakrętki do krawędzi po przeciwnej stronie od kanału. Zbędne linie kanału przycięto. Wyciąganie kształtu można przeprowadzić na dwa sposoby: symetryczny i asymetryczny. Metoda symetryczna jest dość prosta: wybieramy wszystkie elementy szkicu z wyjątkiem otworu, włączamy wyciąganie, jako kierunek wybieramy „Symetrycznie”. Jako wysokość wyciągnięcia podaje się połowę pożądaną grubości elementu, w naszym przypadku element ma mieć 6 mm, więc wymiar powinien wynosić 3 mm. Następnie wykonuje się drugie wyciągnięcie

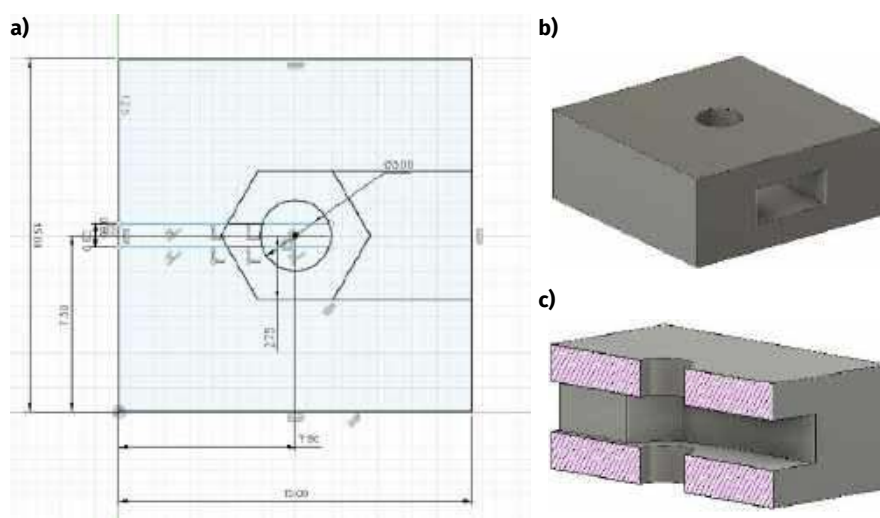
Tabela 1. Wybrane wymiary standardu ISO dla wkrętów i nakrętek z gwintem metrycznym

Rozmiar	Średnica gwintu [mm]	Średnica trzpienia [mm]	Średnica łba guzikowego lub płaskiego [mm]	Średnica łba i nakrętki sześciokątnych [mm]	Maksymalna szerokość łba lub nakrętki [mm]	Grubość nakrętki [mm]
M2	2	1,6	2,6	4	4,7	1,6
M2,5	2,5	2,05	3,1	5	5,8	2
M3	3	2,5	3,6	5,5	6,4	2,4
M4	4	3,3	4,8	7	8,1	3,2
M5	5	4,2	5,8	8	9,4	4
M6	6	5	7	10	11,5	5
M8	8	6,8	10	13	15	6,5
M10	10	8,5	12	16/17	19,6	8



Rysunek 1. Montaż nakrętki we wgłębieniu: projekt (a) i gotowy model (b)

Fotografia 1. Gotowy model z rysunku 1



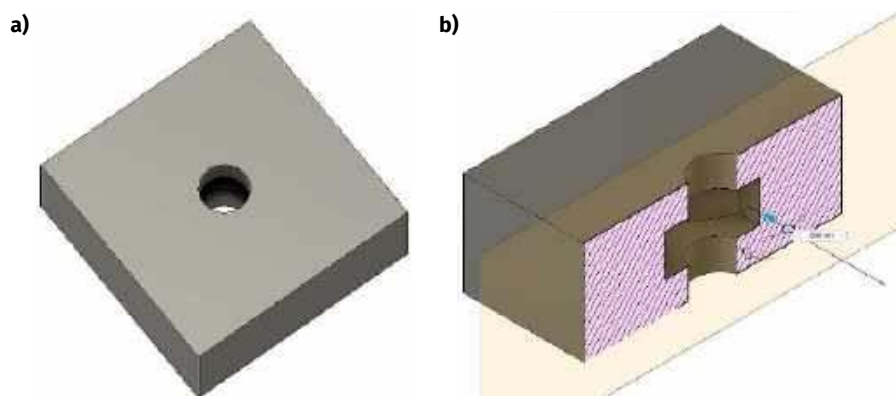
Rysunek 2. Montaż nakrętki wewnątrz kanału: projekt (a), gotowy model (b) i jego przekrój (c)

Fotografia 2. Gotowy model z rysunku 2

wybierając kanał, sam kształt nakrętki i szerszy kanał po drugiej stronie. Tym razem wysokość symetrycznego wyciągnięcia to połowa grubości nakrętki, czyli 1,2 mm. Metoda asymetryczna jest nieco bardziej skomplikowana. Najpierw wyciągamy cały element na pożądaną grubość, czyli 6 mm. Następnie wybiera się oba kanały i kształt nakrętki, i wyciąga się je na wysokość 6...1,8 mm, tj. od grubości całkowitej odejmujemy grubość elementu minus grubość nakrętki, podzielone przez dwa. Trzecim krokiem jest kolejne wyciągnięcie tych samych fragmentów szkicu na wysokość 1,8 mm – należy się przy tym upewnić, iż wybrana jest operacja „Połączenie”. **Rysunek 2b** pokazuje gotowy element, a **rysunek 2c** jego przekrój uzyskany za pomocą narzędzia „Analiza przekroju”

z sekcji „Sprawdź”. Jak wspomniano, to rozwiązanie ukrywa nakrętkę wewnątrz elementu, ale dzięki dodaniu kanałów nakrętka ta może zostać wyjęta lub wymieniona. Większa liczba połączeń poziomych warstw z zewnętrznymi krawędziami elementu ułatwia przenoszenie naprężeń. Jedyną wadą tego rozwiązania jest pogorszona estetyka. **Fotografia 2** prezentuje gotowy element z nakrętką.

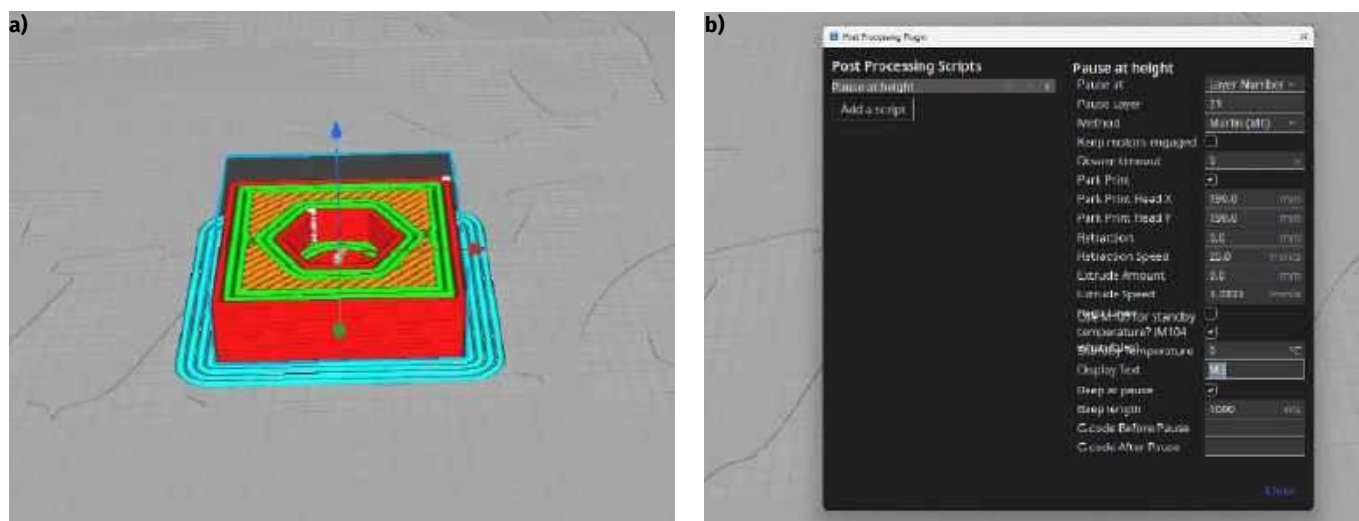
Trzecie rozwiązanie, tym razem z ukrytą nakrętką, wykonuje się analogicznie do rozwiązania drugiego, ale bez kanałów. **Rysunek 3a** pokazuje taki element, a **rysunek 3b** jego przekrój. Nakrętka umieszczana jest w elemencie w trakcie drukowania. Wymaga to jednak modyfikacji programu G-Code i wprowadzenia pauzy w procesie drukowania, w czasie której nakrętka (jedna lub kilka na tej samej



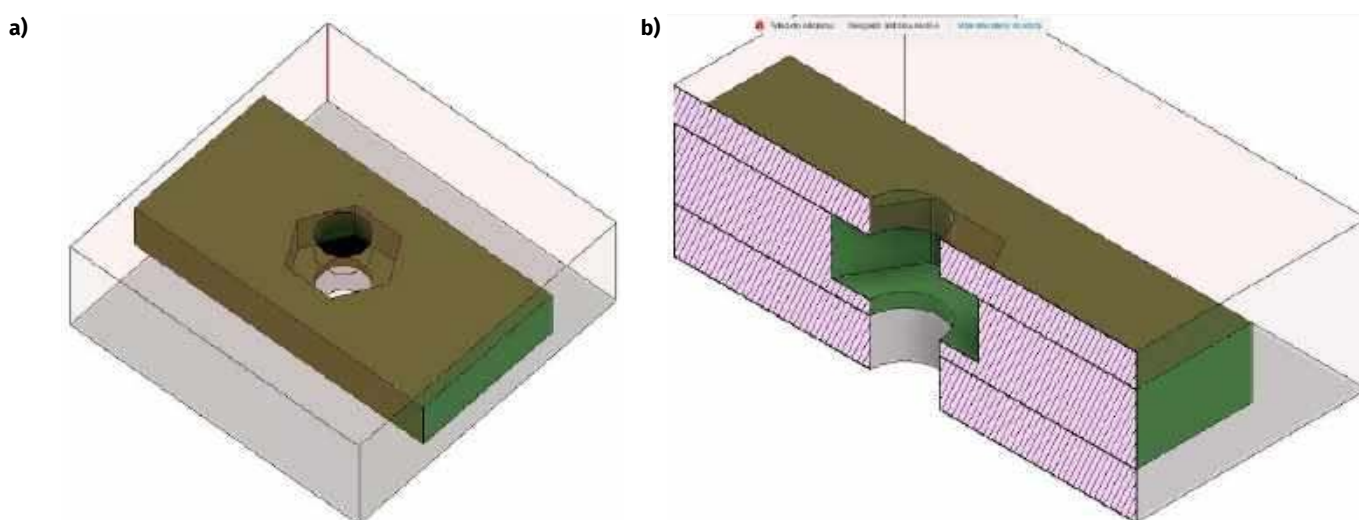
Rysunek 3. Nakrętka ukryta wewnątrz wydruku: model (a) i jego przekrój (b)



Fotografia 3. Gotowy model z rysunku 3



Rysunek 4. Szukanie ostatniej warstwy z otworem na nakrętkę w programie Cura (a) i dodawanie polecenia pauzy (b)

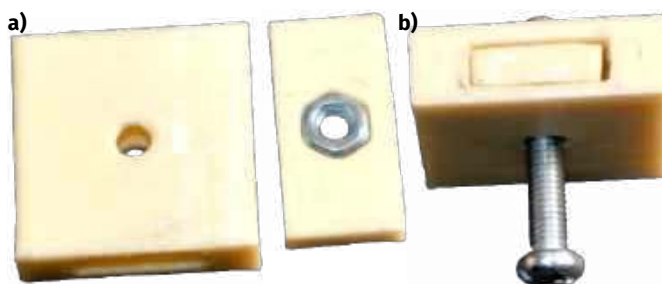


Rysunek 5. Nakrętka umieszczona w oddzielnej płytce, która wzmocni konstrukcję: model (a) i przekrój (b)

głębokości) jest umieszczana na miejscu. W ten sposób można efektywnie zastąpić wtapiane wkładki gwintowane, a także tworzyć pokrętła do ścisków i imadeł. W razie uszkodzenia nakrętki jedyną opcją pozostaje wydrukowanie nowego elementu. Metodę tę można też stosować do ukrywania magnesów neodymowych i innych obiektów. Autor wykonał w ten sposób m.in. pusty w środku brelok, w którego wnętrzu kryją się stalowe kulki. Proces modyfikacji generowanego programu przebiega następująco:

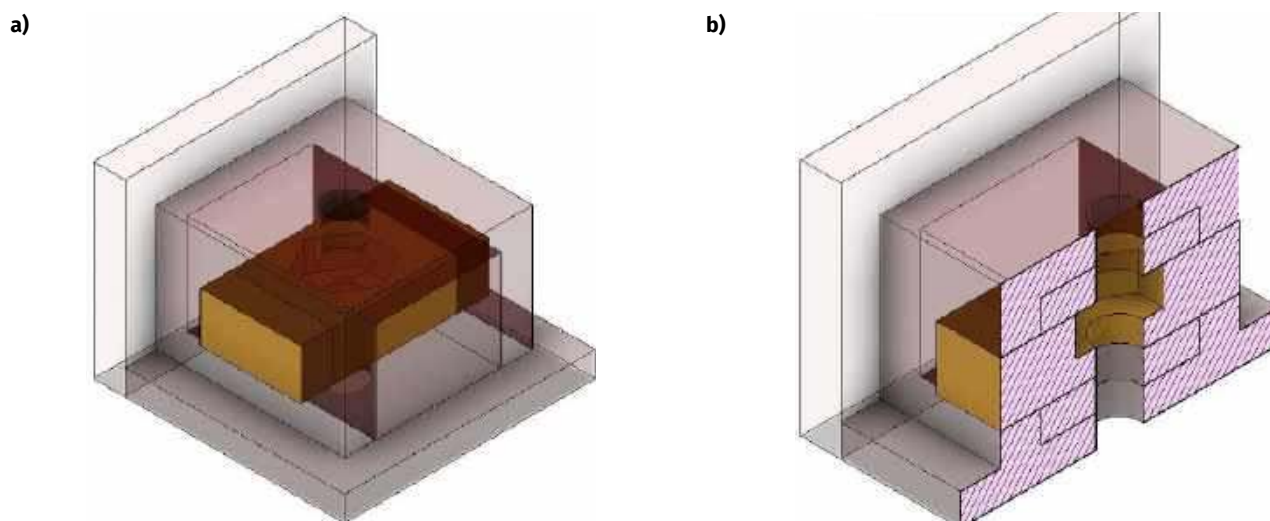
1. W slicerze generuje się normalną sekwencję wydruku.
2. Następnie wybiera się podgląd wydruku.
3. Używając stosownego suwaka, należy znaleźć ostatnią warstwę, w której widać otwór na nakrętkę (rysunek 4a). W tym przypadku jest to warstwa 21.
4. Kolejnym krokiem jest dodanie skryptu pauzy. W Cura należy wybrać Extensions > Post-processing > Modify G-Code.
5. W nowym oknie wybiera się Add Script, a następnie Pause at height.
6. Jako warstwę podaje się warstwę wyżej, należy też zaznaczyć opcję blokady silników oraz podać przesunięcie głowicy (rysunek 4b). Warto dodać też opcję Beep, dzięki której drukarka powiadomi nas o osiągnięciu żądanej wysokości.
7. Ostatnim krokiem jest wygenerowanie sekwencji od nowa i zapisanie pliku G-Code.

Po zakończeniu drukowania warstwy 21 drukarka zatrzyma się i wyda sygnał dźwiękowy. Należy wtedy w odpowiednie miejsce wcisnąć nakrętkę, upewnić się, że jest równo z warstwą, a następnie kontynuować drukowanie.

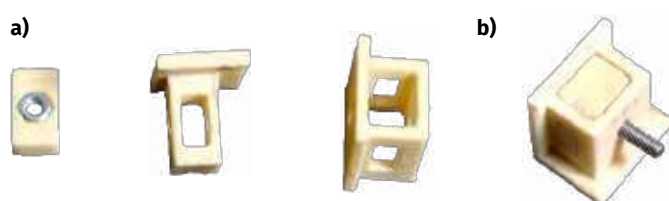


Fotografia 4. Gotowy model z rysunku 5 przed złożeniem (a) i po złożeniu (b). Należy zwrócić uwagę na kierunek warstw w głównej części modelu – są prostopadłe do płytki z nakrętką

Kolejna metoda łączy elementy metody pierwszej i drugiej oraz wymaga dwóch elementów. Pierwszym z nich jest prosty kształt: kwadrat lub prostokąt z otworem na nakrętkę. Może mieć on grubość nakrętki lub być nieco grubszy, jak w metodzie pierwszej. Właściwa część obudowy ma otwór dopasowany wielkością do „wkładki” z nakrętką. Przykładowy model i jego przekrój pokazują rysunki 5a i 5b, zaś fotografia 4 prezentuje gotowy komplet elementów. Ponieważ między warstwami wydrukowanej części mamy mniejszą wytrzymałość, taka wkładka – wydrukowana płasko na blacie – może zwiększyć wytrzymałość obudowy, zwłaszcza jeśli będzie przebiegać przez niemal całą jej wysokość. Nie ma też potrzeby, by wkładka była wciskana w otwór – zamiast tego można zaprojektować płytkę z zagłębieniem na nią, w które zostanie ona wklejona. Technika



Rysunek 6. Dwie części modelu łączone ze sobą za pomocą wkładki. Wkręt i nakrętka nie przenoszą obciążeń. Model (a) i jego przekrój (b)



Fotografia 5. Model z rysunku 6 przed złożeniem (a) i po złożeniu (b). W celu dopasowania elementy były szlifowane. Należy zwrócić uwagę na kierunek warstw – zawsze prostopadle do otworu, jeśli mocowanie ma przenosić obciążenia

tę można też uzupełnić o metodę opisaną wyżej i ukryć wkładki w ściankach obudowy, pauzując drukarkę przed nałożeniem górnych, litych warstw. Nic też nie stoi na przeszkodzie, by wkładki były metalowe, w formie płytek, prętów czy choćby mosiężnych, gwintowanych kołków dystansowych, stosowanych do montażu płytek drukowanych w pewnej odległości od obudowy lub do montażu jednej płytki nad drugą.

Łącząc wymienione dotychczas techniki, możemy pójść o krok dalej, do rozwiązania, w którym wkręt metryczny i nakrętka nie przenoszą obciążeń, a jedynie trzymają właściwy element blokujący na miejscu. Rozwiązanie to pokazują **rysunki 6a i 6b** oraz **fotografia 5**. Elementy są trzymane razem za pomocą wkładki z nakrętką. W tym przykładzie zastosowałem wkładkę z nakrętką wsuwaną, pokazaną już wcześniej. Element ten jest umieszczany w specjalnie przygotowanym otworze w obu elementach tak, by te nie mogły się przemieścić względem siebie. Wkręt blokuje położenie wkładki. Wytrzymałość połączenia zależy od wytrzymałości mechanicznej wydrukowanych elementów. Warto też pamiętać, że pokazane rozwiązania mogą wymagać dodatkowej obróbki mechanicznej lub/i uwzględnienia tolerancji w projekcie – inaczej elementy mogą do siebie zwyczajnie nie pasować. Ostatnie dwa przykłady pokazane na fotografiach wymagały dodatkowego szlifowania pilnikiem.

Wtapiane wkładki gwintowane oraz gwintowanie plastiku

W projekcie obudowy opisanej w poprzednim artykule całkowicie zrezygnowano z nakrętek, opierając się tylko i wyłącznie na tworzeniu „naturalnego” gwintu podczas pierwszego skręcania

całości. Rozwiązanie to ma zaletę w postaci niezwyklej wręcz prostej – wystarczy zaprojektować otwór o średnicy rdzenia wkrętu. Pod względem wytrzymałości rozwiązanie to jest niemal tak samo dobre, jak użycie nakrętki, szczególnie gdy długość gwintu jest kilka razy większa od grubości nakrętki. Do wad należy fakt, iż po kilku lub kilkunastu użyciach gwint zwyczajnie się wyrobi i wkręt nie będzie się dobrze trzymał. Jest to zatem rozwiązanie generalnie jednorazowe. Zamiast standardowych wkrętów metrycznych można zastosować blachowkręty lub wkręty do plastiku. Te mają większą średnicę i mniejszą gęstość gwintu, dzięki czemu lepiej (głębiej) wcinają się w tworzywo. Połączenia takie mogą być pewniejsze, ale też nie będą wieczne.

Najlepsze rozwiązanie stanowi zastosowanie wtapianych wkładek gwintowanych. Te drobne elementy wykonane z mosiądzu mają wewnątrz gwint metryczny, a na zewnątrz są uformowane tak, by tworzyć dużą powierzchnię styku z tworzywem, a przy tym nie dać się łatwo wyrwać lub wykręcić. Instalacja insertów podnosi koszty produkcji, dlatego stosuje się ją głównie w produktach komercyjnych wyższej jakości (lepszych obudowach uniwersalnych, laptopach itp.). W przypadku małych serii i pojedynczych sztuk wykonywanych metodą druku 3D wkładki takie mają znacznie większy sens ekonomiczny – szczególnie hobbysty umiłowali sobie to rozwiązanie w ostatnich latach, a na rynku dostępne są prasy z elementem grzejącym do ich instalowania. Jednakże takie wkładki można z powodzeniem zamontować także za pomocą zwykłej kolby lutowniczej. Należy zaprojektować otwór pod wkładkę wedle specyfikacji producenta. **Tabela 2** pokazuje przykładowe zalecenia.

W następnym kroku należy umieścić nad otworem wkładkę i docisnąć grotiem lutownicy (są do tego celu dostępne specjalne, walcowe groty, ale zwykły grot stożkowy lub płaski też się sprawdzi przy odrobinie doświadczenia) o temperaturze 200...300°C, zależnie od rodzaju tworzywa. Rozwiązanie to ma też taką zaletę, że wokół wkładki kolejne warstwy tworzywa są ze sobą lepiej stapiane, co dodatkowo zwiększa wytrzymałość połączenia.

W niektórych przypadkach, szczególnie gdy planujemy użyć filamentu mniej kruchego niż PLA, można rozważyć zrezygnowanie ze śrub i zastąpienie ich zintegrowanymi klipsami lub dodatkowymi spinakami. Takie rozwiązanie zostanie pokazane w przyszłości.

Tabela 2. Przykładowe wymiary typowych wtapianych wkładek gwintowanych

	M2	M2,5	M3	M4	M5	M6
Średnica wkładki	3,5 mm	3,5 mm	5 mm	6 mm	7 mm	9 mm
Średnica otworu (projektowa)	3,2 mm	3,2 mm	4,6 mm	5,4 mm	6,6 mm	8,7 mm

Zagadnienia mechaniczne i termiczne projektowania obudowy do druku 3D

Projektując obudowę, która będzie wykonywana na filamentowej drukarce 3D, należy zawsze uwzględniać aspekty mechaniczne i termiczne takiego projektu. Mechaniczne – ze względu na to, że większość dostępnych filamentów ma znacząco mniejszą wytrzymałość między warstwami. Obudowy wykonane z tworzyw termoplastycznych są też podatne na deformację pod wpływem temperatury, dlatego aspekt termiczny należy uwzględnić w projekcie. W większości przypadków wyklucza to więc użycie PLA jako materiału do wykonywania obudowy. Materiał ten bowiem może zacząć się deformować pod obciążeniem już w temperaturze 50°C. W temperaturze 80°C staje się zbyt plastyczny i deformuje się nawet pod własnym ciężarem. Użycie bardziej wytrzymałego materiału, jakim jest PETG, nie zwalnia jednak konstruktora z wymogu zwrócenia uwagi na własności termiczne tworzywa.

Obudowa, która będzie tu prezentowana, przeznaczona jest do zasilacza regulowanego. W związku z tym wewnątrz znajdą się dwa znaczące źródła ciepła: radiator regulatora oraz zasilacz sieciowy. By uniknąć sytuacji, w której po dłuższym czasie pracy obudowa zaczyna wyglądać jak z obrazu Salvadora Dali, a przy tym by ograniczyć wpływ temperatury na stabilność zasilacza, projekt z góry zakładał wymuszone chłodzenie za pomocą wentylatora. Zaawansowane programy CAD, jak na przykład SolidWorks czy CATIA, pozwalają na przeprowadzenie symulacji przepływu powietrza wewnątrz obudowy. My jednak nie będziemy korzystać z takich narzędzi, gdyż w tym wypadku zwyczajnie nie jest to konieczne. Kształt wnętrza będzie nieco bardziej skomplikowany, by wymusić obieg powietrza w pożądanym kierunku, ale te dodatkowe elementy posłużą też do mechanicznego wzmocnienia całej konstrukcji. Zwiększy to też koszt materiałowy i wydłuży czas wydruku, ale na to niewiele można poradzić.

Tworząc projekt obudowy, trzeba wstępnie przyjąć, jaką będzie ona miała formę i jakie są minimalne wymagania względem jej wielkości. Parametry te zależą od wielkości radiatora, zasilacza sieciowego i wentylatora. Należy też przyjąć z góry, w jakiej pozycji i jaką metodą będzie montowany radiator oraz czy zamierzamy użyć dodatkowych elementów, by powietrze przepływało przez wszystkie jego żeberka, czy raczej skupimy nawiew na środkowej części radiatora, dodając otwory wentylacyjne nad i pod radiatorem, by naturalna konwekcja chłodziła jego skraje. Można też dodać więcej wentylatorów, by mieć lepsze pokrycie całego radiatora.

Kolejną kwestią jest instalacja różnych złączy w obudowie: gniazda IEC do kabla zasilającego i złączy bananowych do przewodów wyjściowych zasilacza. Elementy te będą poddane częstym obciążeniom w trakcie użytkowania, więc obudowa w tych miejscach powinna być wzmocniona. Podobnie jest w przypadku panelu kontrolnego, który powinien być wykonany nieco solidniej, by się nie ugiął w trakcie użytkowania. Do płyty czołowej przykręcona będzie płytka z wyświetlaczem, przyciskami i enkoderem. Od wewnętrznej strony można zatem dodać zgrubienia, które usztywnią całość. Warto przy okazji zastanowić się nad metodą montażu panelu do reszty obudowy. Najprościej by było go zwyczajnie przykręcić. Panel stanie się wtedy kolejnym łącznikiem między bokami obudowy, zwiększając jej sztywność. Same wkręty można ukryć wewnątrz obudowy, by poprawić estetykę, ale kosztem łatwości demontażu w razie konieczności naprawy. Za to umieszczenie wkrętów na zewnętrznej powierzchni panelu lub na ściankach bocznych znacząco ułatwia montaż i demontaż, ale pogarsza estetykę całości. Przejdźmy zatem do tego aspektu projektowania.

Zagadnienia projektowania estetycznych i ergonomicznych obudów

Autor pragnie zaznaczyć, że nie jest specjalistą od wzornictwa przemysłowego. Jest to pewne utrudnienie, ale podstawowe zasady estetyki i ergonomii są dość uniwersalne, a przy tym łatwe do zastosowania. O ile o gustach z zasady się nie dyskutuje, o tyle każdy się

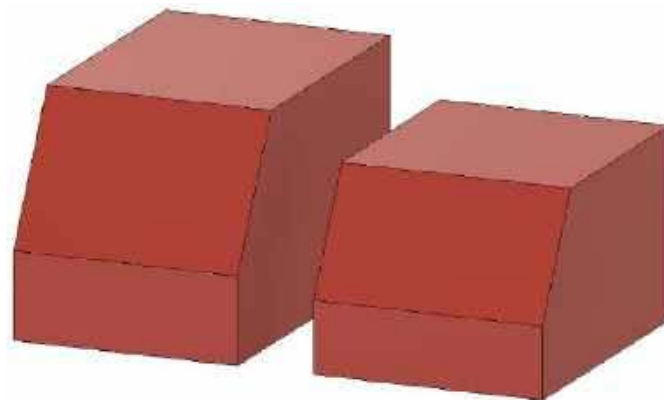
zgodzi, że, na przykład, kolor czarny pasuje w elektronice niemal do wszystkiego.

Obudowa powinna być proporcjonalna i mieć raczej prostą bryłę. Druk 3D pozwala na tworzenie dość fantastycznych, złożonych kształtów, ale my projektujemy obudowę urządzenia do warsztatu, a nie małe dzieło sztuki – stąd prosta, niemal brutalistyczna, użyteczna forma. Obowiązują nas też trzy ograniczenia:

1. minimalne wymiary komponentów i płytek drukowanych,
2. maksymalne wymiary pola roboczego,
3. właściwości filamentu.

Nie bierzemy pod uwagę czasu drukowania ani liczby elementów, z których będzie wykonana obudowa – to nie produkcja wielkoseryjna, gdzie liczy się każdy gram i każda sekunda. Zresztą w przypadku produkcji wielkoseryjnej zakłada się, że obudowa tworzywowa będzie wykonywana na wtryskarce, nie na drukarce 3D. W tym przypadku wytrzymałość materiału jest taka sama w każdej płaszczyźnie, ale zbyt skomplikowane kształty trzeba wykonać z większej liczby elementów lub z nich całkowicie zrezygnować. Najprościej zrozumieć tę różnicę, patrząc na modele przeznaczone do druku 3D z takich stron, jak na przykład Printables (<https://www.printables.com>) czy Thingiverse (<https://www.thingiverse.com>) i porównując je z zabawkami, obudowami fabrycznych urządzeń czy z przedmiotami codziennego użytku. Standardem w przypadku urządzeń jest obudowa składająca z dwóch części skręconych lub spiętych razem, przy czym projektant bardzo się postarał, by nie było żadnych wypukłości czy wgłębień, które wymagałyby formy składającej się z więcej niż dwóch elementów. Jeśli kształt jest bardziej złożony, to najczęściej wykonany jest z większej liczby elementów, gdyż dwie lub trzy proste formy wtryskowe są dużo tańsze od jednej, wieloczęściowej formy z dodatkowymi suwakami.

Wracając jednak do naszego projektu obudowy, głównym limitem wymiarów jest pole robocze drukarki 3D. Drukarka Ender 3 V2 ma pole 220×200×250 mm. W praktyce jednak nie powinno przekraczać się 200×200 mm w poziomie i 240 mm w pionie, by mieć jakieś marginesy. W drugą stronę ograniczają nas przede wszystkim rozmiar zasilacza sieciowego, radiatora i wentylatora. Minimalny wymiar panelu przedniego to długość płytki drukowanej – jej szerokość jest mniejsza od szerokości radiatora. Jeśli ustawimy radiator bokiem, wtedy możemy zmniejszyć szerokość obudowy do szerokości płytki. Na **rysunku 7** pokazano zgrubne bryły zasilacza. W obu doliczono marginesy na grubość ścianek, ale w jednej szerokość zależy od szerokości radiatora, a w drugiej – od szerokości płytki panelu kontrolnego. Płytkę panelu ma szerokość 89 mm i długość 53 mm. Radiator ma wymiary 125×72×35 mm. Standardowy wentylator 80×80 mm ma grubość 25,5 mm, ale zostawimy dla niego dodatkowe 4,5 mm przestrzeni. Modułowy zasilacz impulsowy ma wymiary 85×50×33 mm. W obu



Rysunek 7. Przykładowe bryły obudowy zasilacza: po lewej obudowa węższa, ale głębsza i wyższa. Po prawej obudowa szersza, ale przez to niższa i płytsza. W obu zastosowano zasadę złotego podziału dla poprawy estetyki

przypadkach założono nachylony pod kątem 30 stopni panel kontrolny oraz identyczną wysokość. By nieco poprawić estetykę brył, zastosowano też **zasadę złotego podziału**. Najprościej jest przyjąć, iż **stosunek dwóch wymiarów (np. długości) powinien mieć wartość 1:1,618**. Jeśli zatem obudowa ma wysokość 120 mm, to dolna część frontu powinna mieć wysokość 45,8 mm.

Obie bryły wyglądają dobrze i mają zbliżone wymiary: obudowa po lewej ma wymiary 112×112×180 mm, po prawej zaś: 93×114,5×150 mm. O wyborze jednej z nich będą decydować raczej aspekty konstrukcyjne, tj. to, jak łatwo da się w niej upakować wszystkie elementy. Węższa bryła będzie lepiej pasować do innych, typowych elementów warsztatu elektronika: komercyjnych zasilaczy, stacji lutowniczych czy multimetrów. Z kolei szersza obudowa lepiej nadaje się do generatorów funkcyjnych, multimetrów stołowych czy oscyloskopów. Oczywiście nic nie stoi na przeszkodzie, by zaprojektować oba warianty i dać wybór docelowemu użytkownikowi. Autor skłania się jednak w stronę węższej obudowy, gdyż jej konstrukcja ułatwi realizację dobrego chłodzenia z wymuszonym obiegiem powietrza.

Panel przedni zawiera wyświetlacz, enkoder, cztery przyciski i dwie diody sygnalizacyjne. Na płycie znajdują się przyciski typu tact – w amatorskich konstrukcjach zazwyczaj pozwala się im wystawać z obudowy przez okrągłe otwory. Jest to jednak rozwiązanie nieestetyczne i niegodne profesjonalisty. Zamiast tego każdy przycisk otrzyma nakładkę ze stosownym oznaczeniem funkcji na zewnętrznej powierzchni. Przy produkcji wielkoseryjnej takie oznaczenia nanoszone są ze wzornika za pomocą specjalnego gumowego tamponu. Nam jednak będzie dużo łatwiej zaprojektować

wgłębienia w kształcie pożądanego symbolu, a potem wypełnić je farbą akrylową lub inną, w pożądanym kolorach. Podobnie z oznaczeniami na samym panelu przednim. W przypadku enkodera można zaprojektować własny kształt gałki, dopasowany do formy obudowy.

Jeśli chodzi o wyświetlacz, to powszechnie występującym, karygodnym błędem hobbystów jest wykonywanie otworu wielkości metalowej ramki ekranu LCD, przez co ten w całości wręcz wystaje, a blask podświetlenia jest widoczny dookoła niego. Nie dość, że takie rozwiązanie wygląda zwyczajnie źle, to jeszcze w żaden sposób nie chroni wyświetlacza przed uszkodzeniami mechanicznymi. Poprawną metodą jest uczynienie otworu niewiele większym od samego pola znakowego, dzięki czemu nie widzimy nawet krawędzi okienka wyświetlacza, nie wspominając o metalowej ramce. Od wewnętrznej strony obudowy metalową ramkę można otoczyć rantem z tworzywa, który dodatkowo wzmocni tę część panelu. Tak też postąpimy, posługując się notą katalogową wyświetlacza. Dokumentacje techniczne pozwolą nam również dopasować otwory pod gniazdo IEC i włącznik zasilania.

Zakończenie

Omówiliśmy pobieżnie kilka metod składania obudów i innych elementów wykonanych metodą filamentowego druku 3D (FDM). Przyjrzelśmy się też zagadnieniom konstrukcyjnym i estetycznym naszej przykładowej obudowy. W następnej części przejdziemy do faktycznego projektu oraz omówimy zastosowane rozwiązania. Na koniec zaś zobaczymy, jak się prezentuje gotowe urządzenie.

Paweł Kowalczyk, EP

REKLAMA

Wydawnictwo AVT nawiąże współpracę redakcyjną z osobami dobrze operującymi terminologią elektroniki i słowem pisanym. Propozycja szczególnie interesująca dla nauczycieli elektroniki, autorów artykułów, skryptów i książek.

**Aplikacje prosimy kierować na adres:
redakcja@elportal.pl**



Internet Rzeczy w pomiarach środowiskowych (13)

Dołączanie rezystancyjnych i pojemnościowych czujników deszczu do modułu Enviro Weather

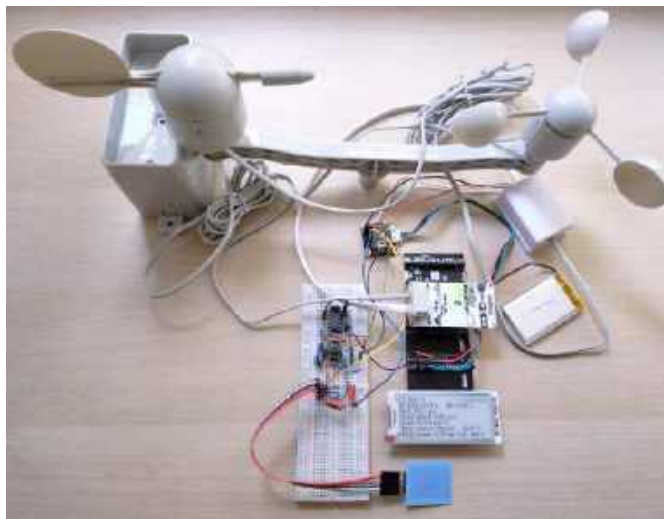
Czujnik opadów deszczu może znacznie zwiększyć możliwości monitorowania parametrów środowiska, zwłaszcza w erze zmian klimatycznych. Pojemnościowy sensor RC-SPC1K firmy Radiocontrolli pozwala na dokładne pomiary poziomu opadów przy zastosowaniu stosunkowo prostego oraz taniego układu pomiarowego i doskonale uzupełnia stację pogodową Enviro Weather.

W czujnikach deszczu stosowane są różne sposoby pomiaru, np.: rezystancyjne, optyczne, pojemnościowe i akustyczne. Jedną z najprostszych metod wykrywania opadów jest zastosowanie rezystancyjnego czujnika grzebieniowego, opartego na dwóch układach przewodzących ścieżek, umieszczonych w niewielkiej odległości. Kształt zachodzących na siebie grzebieni umożliwia wykrycie nawet niewielkich, pojedynczych kropeł deszczu, które – zwierając obydwa grzebienie – powodują spadek rezystancji czujnika.

Najczęstszą praktyką w wykrywaniu kropeł deszczu na przedniej szybie samochodu jest wykrywanie światła podczerwonego przechodzącego przez wewnętrzne ścianki szyby przedniej. Gdy na zewnętrznej powierzchni szkła znajdują się krople deszczu, następuje załamanie światła (rozpraszanie wiązki), które redukuje ilość światła powracającego do czujnika w porównaniu z pierwotnymi warunkami (bez obecności wody). Technika ta opiera się na statystycznym prawdopodobieństwie, że gdy pada deszcz, na powierzchni szkła znajdują się krople deszczu przecinające drogę wiązki [8].

Technologia pomiaru pojemnościowego opiera się na wykrywaniu zmian pojemności w wyniku modulacji stałej dielektrycznej materiału oddzielającego przewodniki elektryczne (okładki) kondensatora przez krople wody, które znajdują się na powierzchni czujnika. Pojemnościowy czujnik deszczu precyzyjnie i skutecznie wykrywa tempo, aktualną intensywność oraz moment zakończenia opadów, a ponadto pozostaje nieczuły na fałszywe alarmy spowodowane zanieczyszczeniami, które w czujnikach optycznych powodowałyby odbicia imitujące odbicie deszczu. Ponadto, podczas gdy rezystancyjne czujniki deszczu są wrażliwe na korozję i zanieczyszczenia, topologia pojemnościowa nie ma tych wad [9].

Czujnik akustyczny składa się z okrągłej pokrywy ze stali nierdzewnej, zamontowanej na sztywnej ramie. Pod pokrywą znajduje się detektor piezoelektryczny. Krople deszczu uderzają w powierzchnię czujnika z prędkością, która jest funkcją średnicy



Autor składa podziękowania panu Maciejowi Michnie z Centrum badań i rozwoju Nordic Semiconductor w Krakowie za udostępnienie zestawów sprzętowych Power Profiler Kit II (PPK2).



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

kropki deszczu. Pomiar intensywności opadów opiera się na akustycznej detekcji każdej pojedynczej kropli deszczu uderzającej w pokrywą czujnika. Większe krople generują silniejszy sygnał akustyczny niż mniejsze. Detektor piezoelektryczny zamienia sygnały akustyczne na napięcie. Wynik pomiaru obliczany jest na podstawie sumy indywidualnych sygnałów napięciowych na jednostkę czasu, przy znanej powierzchni czujnika. Ponadto można obliczyć intensywność i czas trwania deszczu [7].

Rezystancyjny czujnik deszczu Steam Sensor

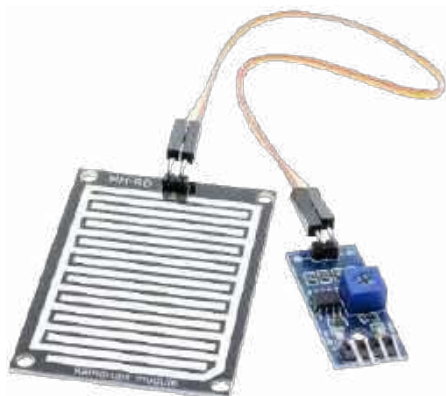
Czujnik Steam Sensor KS0203 firmy Keyestudio jest analogowym czujnikiem rezystancyjnym (**fotografia 1**) o napięciu zasilania 3,3...5 V. Dwie elektrody w postaci gwiazdki ułożone są naprzemiennie – jedna z nich jest dołączona do masy, a druga – przez rezystor 1 MΩ – do zasilania. Do tej elektrody dołączone jest wyjście analogowe przez rezystor 470 Ω. Czujnik wykrywa krople deszczu oraz parę. Napięcie wyjściowe wzrasta wraz z wilgotnością na powierzchni PCB. Czujnik jest bardzo tani (kilka złotych) i łatwo dostępny [10].

Rezystancyjny czujnik deszczu Raindrops Module

Raindrops Module składa się z dwóch części: sondy pomiarowej i modułu detektora, połączonych przewodami (**fotografia 2**). Sonda jest dołączona do masy oraz – przez rezystor 10 kΩ – do zasilania VCC. Sygnał z grzebieni jest wyprowadzony na wyjście analogowe



Fotografia 1. Czujnik deszczu i pary Steam Sensor [10]



Fotografia 2. Czujnik Raindrops Module [11]

A0, a jednocześnie podawany na wejście komparatora LM393 (zasilanie 2...30 V, 1 mA). Wyjście komparatora pozostaje dołączone do wyprowadzenia D0 modułu. Próg zadziałania komparatora można ustawić potencjometrem. Gdy sonda jest sucha, na wyjściu analogowym A0 panuje napięcie zbliżone do VCC. W miarę wzrostu wilgotności sondy napięcie A0 maleje. W artykule [11] pokazano odpowiedź sondy na krokowy wzrost zawilgocenia. Czujnik jest tani i oferowany przez wiele firm, w sprzedaży znajduje się też wersja z tym samym modulem detektora i dołączonym czujnikiem wilgotności gleby (YL-69).

Pojemnościowy czujnik deszczu RC-SPC1K

Czujnik opadów RC-SPC1K firmy Radiocontrolli jest wykonany w technologii grubowarstwowej na podłożu allumina (Al_2O_3) [1]. Jest to materiał o dużej niezawodności elektryczno-termicznej. Czujnik składa się z trzech elementów: sensora pojemnościowego (strona A) oraz grzałki i czujnika temperatury (strona B) – patrz **fotografia 3**.

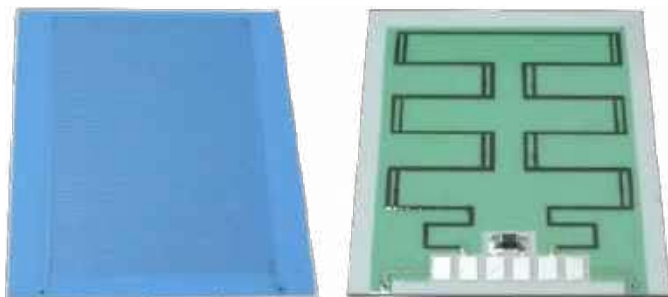
Podstawowe parametry czujnika RC-SPC1K [1]:

- Pojemność początkowa (suchy): 100 pF $\pm 10\%$
- Rezystancja grzałki: 42 Ω $\pm 10\%$
- Czujnik temperatury NTC (25°C): 1 k Ω $\pm 10\%$
- Wymiary: 30,48×35,56 mm

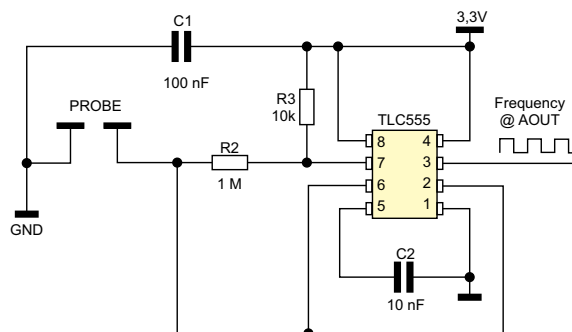
Dostępna jest także wersja RC-SPC1KA o większej zmianie pojemności w funkcji wilgotności (ponad 3000 pF @ 100%).

Powierzchnia strony A (niebieska) to obszar wrażliwy (czujnik pojemnościowy), wystawiany na działanie deszczu. W przypadku opadów pojemność wzrasta do wartości wyższej niż w warunkach suchych (ponad 420%), do 550 pF przy 100% wilgotności. Z tyłu czujnika deszczu znajduje się czujnik temperatury NTC oraz grzejnik rezystancyjny. Dzięki przewodności cieplnej ceramiki ciepło emitowane przez grzałkę z tyłu czujnika jest natychmiast przenoszone na górną powierzchnię. Grzejnik zapewnia szybkie wysychanie powierzchni detekcji, chroniąc powierzchnię przed mgłą, skroploną wilgocią i szronem.

W dokumentacji firmowej proponowane są dwa sposoby detekcji zmian pojemności czujnika: włączenie go jako element generatora astabilnego (na jednej bramce logicznej) lub dokonywanie pomiaru czasu odpowiedzi na pobudzenie impulsem prostokątnym. W celu



Fotografia 3. Czujnik deszczu RC-SPC1K (strona A – po lewej) [1]



Rysunek 1. Schemat układu do pomiaru opadów deszczu [2]

pomiaru temperatury wystarczy dołączyć rezystor 1 k Ω do NTC i mierzyć zmianę napięcia na dzielniku. Wysterowanie grzałki wymaga zastosowania tranzystora kluczującego.

Układ pomiarowy opadów deszczu

W celu uzyskania stabilnej i dokładnej pracy czujnika zastosowano układ generatora astabilnego z popularnym układem NE555 w wersji MOS (**rysunek 1**). Układ został zaadaptowany z konstrukcji miernika wilgotności gleby [2]. Do pojemności czujnika wynoszącej 100 pF został dobrany rezystor R2 tak, aby częstotliwość dla suchego czujnika wynosiła 167 kHz:

$$F_{gen} = 1,44 / ((R3 + 2 \cdot R2) \cdot C_p)$$

przy $C_p = 100$ pF i $R2 = 22$ k Ω został użyty $R3 = 10$ k Ω .

Dodatkowo zastosowano dzielnik cyfrowy SN74HC4060 z podziałem dwunastobitowym (wyprowadzenie Q3). Uzyskany przebieg prostokątny ma częstotliwość:

$$F_{out} = F_{gen} / 4096$$

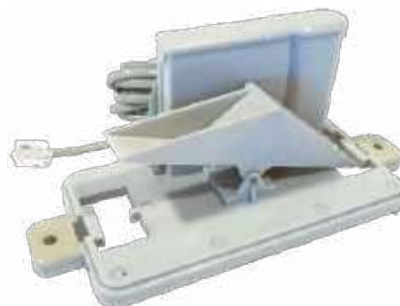
Maksymalna częstotliwość czujnika w suchym powietrzu wynosi: $F_{gen} = 167,2$ kHz, $F_{out} = 40,65$ Hz, a minimalna (czujnik w wodzie): $F_{gen} = 41,26$ kHz, $F_{out} = 10,75$ Hz.

Do stabilnej pracy układu potrzebne jest czyste napięcie zasilania pozbawione szpilek. Nie nadaje się do tego napięcie z szyny 3V3 płytki Pico W, dlatego układ jest zasilany czystym napięciem 3,3 V z płytki ADS1115.

Do termistora NTC czujnika opadów został dołączony rezystor 1 k Ω , podłączony do czystego napięcia 3,3 V. Do tego samego napięcia został dołączony też dzielnik z dwoma rezystorami 3,61 k Ω 0,25% dający napięcie referencyjne. Oba potencjały dołączono do wolnych wejść przetwornika ADS1115 – ich pomiar pozwala na określenie temperatury czujnika opadów i ma istotne znaczenie przy używaniu grzałki czujnika. W dokumentacji firmowej znajduje się tabela wartości rezystancji NTC w funkcji temperatury, występuje jednak problem niskiej dokładności (10%) rezystancji NTC – całość wymaga więc skalibrowania napięcia pomiarowego w znanej temperaturze.

Kubekowy czujnik opadów (COM-B020)

Większość deszczomierzy określa poziom opadów w milimetrach wysokości, w przeliczeniu na 1 m² w określonym czasie [3]. Czujnik deszczu COM-B020, dostarczany z zestawem stacji pogodowej



Fotografia 4. Czujnik opadów po zdjęciu obudowy [3]

Enviro Weather, jest w rzeczywistości prostym urządzeniem mechanicznym – jego działanie opisano już w EP 04/2024.

Modyfikowany moduł DFRobot I²C ADS1115

Moduł DFRobot I²C ADS1115 (DFR0553) firmy DFRobot [12] zawiera układ przetwornika analogowo-cyfrowego ADS1115. Opis samego modułu oraz zalecanej jego modyfikacji znajduje się w EP 10/2024.

Modyfikacja dołączenia czujnika kierunku wiatru do modułu Enviro Weather

Pomiar napięcia wyjściowego dzielnika, złożonego z 8 kontaktów oraz wbudowanych rezystorów czujnika wiatru z zestawu Enviro Weather, jest dokonywany za pomocą przetwornika analogowo-cyfrowego w mikrokontrolerze i pozwala określić kierunek, w który skierowany jest wiatrowskaz. W zestawie Enviro Weather dzielnik jest zasilany z szyny 3V3 RPI Pico W. Jednak zakłócenia szpilkowe na tej szynie praktycznie uniemożliwiają dokładne pomiary napięcia [4].

Napięcie z dzielnika jest podawane na wejście przetwornika ADC procesora RP2040. Wbudowany przetwornik ADC procesora RP2040 jest tylko 12-bitowy. Dla kompatybilności z arytmetyką 16-bitową odczytane słowa danych są przesuwane w lewo o 4 bity. Powoduje to duże skoki wartości dla niewielkich zmian poziomu wejściowego. Brak dokładnego napięcia odniesienia przetwornika powoduje duże skokowe, przypadkowe zmiany odczytu dla impulsowych zakłóceń przetwornicy zasilania procesora. Dlatego w układzie został zastosowany zewnętrzny przetwornik ADC typu ADS1115. Bardzo duża rezystancja wejściowa przetwornika nie zaburza pracy dzielnika napięciowego z rezystorami o wartościach kiloomowych. Zasilanie płytki I²C ADS1115 jest pobierane z szyny VSYS płytki Raspberry Pi PicoW modułu Enviro Weather. Płytka przetwornika ma wbudowany układ LDO typu LP5907MFX-3.3, dostarczający do 250 mA czystego napięcia 3,3 V. Z płytki ADS1115 zostało pobrane czyste napięcie zasilania 3,3 V i dołączone do dzielnika pomiaru kierunku wiatru zestawu COM-B020.

Dołączenie czujników zewnętrznych do modułu Enviro Weather

Do płytki Enviro Weather dołączany jest czujnik opadów deszczu i anemometr zestawu COM-B020 według opisu fabrycznego [3], a czujnik kierunku wiatru – według opisu powyżej. Zasilanie i sygnał pomiarowy rezystancyjnych czujników deszczu Raindrops Module oraz Steam Sensor są dołączone do płytki przetwornika ADS1115. Zasilanie i napięcie dzielnika z termistorem NTC zostały również dołączone do płytki ADS1115. Wyjściowy sygnał prostokątny układu pomiarowego opadów deszczu z czujnika pojemnościowego został natomiast skierowany na wejście GP15 procesora RP2040.

Praca czujników z modułem Enviro Weather

Zastosowanie płytki Enviro Weather wymaga najpierw wpisania do niej najnowszej wersji firmowego pliku obrazu (uf2) zawierającego MikroPythona oraz biblioteki firmowe, np. Pico Graphics. Następnie należy wpisać folder projektu najnowszej aplikacji Enviro. Jest to dokładnie omówione w poprzednim artykule „Stacja pogodowa Enviro Weather firmy Pimoroni” [3].

Z pobranego pliku *Enviro_Rain.zip* (<https://tiny.pl/cdgpcj2j>) podmieniane są pliki *main.py* i *config.py*.

Dokładny opis konfiguracji płytki znajduje się w artykule [3].

Na początku programu wykrywane są wszystkie układy na szynie I²C, a na czarnym tle wyświetlacza pokazywany jest napis „Nowy pomiar”. Następnie inicjalizuje się obsługę Enviro Weather i synchronizuje zegar RTC PCF85063A płytki z czasem pobranym poprzez Wi-Fi z Internetu. Potem następuje aktywacja dostępu do czujników BME280, BME688 i LTR-599 na szynie I²C oraz pobranie



Fotografia 5. Przykład pomiaru bez czujników wiatru i bez opadów



Fotografia 6. Przykład pomiaru podczas niewielkich opadów deszczu i wiatru

wyników pomiarów parametrów z modułu Enviro Weather: temperatury, wilgotności, ciśnienia atmosferycznego, poziomu oświetlenia, prędkości wiatru oraz poziomu i prędkości opadów deszczu.

Następnie wykonywany jest pomiar opadów deszczu na podstawie sygnału z pojemnościowego czujnika RC-SPC1K. Wykrywane są zbocza sygnału i na tej podstawie mierzy się czas półokresu przebiegu prostokątnego (wypełnienie 50%). Wyliczana jest też procentowa pozycja wartości pomiędzy minimalnym (czujnik suchy) i maksymalnym (czujnik mokry) czasem. Na tej podstawie określa się (dosyć arbitralnie) słowną ocenę intensywności opadów. Ostatnim etapem jest odczyt napięcia z czujnika kierunku wiatru i konwersja danych na osiem kierunków świata.

Wszystkie wyniki pomiarów i obliczeń są pokazywane na wyświetlaczu e-paper. W prawym górnym rogu wyświetlana jest data i godzina ostatniego pomiaru, a po lewej: temperatura, wilgotność, ciśnienie atmosferyczne, poziom oświetlenia oraz prędkość wiatru. Ostatni z tych parametrów jest określany na podstawie zliczania impulsów od czujnika. Gdy czujnik nie jest podłączony, to brak impulsów oznacza prędkość zerową (**fotografia 5**). Poniżej wyświetla się kierunek geograficzny wiatru. W przypadku braku czujnika kierunku wiatru na ekranie znajduje się napis *None*. Dalej prezentowane są dane z pojemnościowego czujnika deszczu: w postaci słownej oceny intensywności oraz procentowej intensywności opadów. Na dole znajdują się ilościowe dane z kubelkowego czujnika opadów z zestawu Enviro Weather: pomiar poziomu opadów deszczu i prędkość opadów deszczu (**fotografia 6**).

Na koniec działania aplikacji wywoływana jest firmowa funkcja *enviro.sleep* programująca RTC na wybudzenie procesora oraz wprowadzająca procesor na pewien czas (np. 1 min) w uśpienie. Podczas zasilania z wejścia BAT wyłączane jest zasilanie całej płytki Enviro Weather oraz wszystkich dołączonych czujników, z wyjątkiem układu RTC. Sygnał alarmu z RTC wymusza reset procesora. Oprogramowanie było uruchamiane w środowisku Thonny. Na **listingu 1** jest pokazane okno Shell widoczne po włączeniu aplikacji pomiarowej.

Do dynamicznego pomiaru prądu zasilania bardzo dobrze nadaje się zestaw Power Profiler Kit II (PPK2) firmy Nordic Semiconductor. Pobór prądu całego układu przy zasilaniu bateryjnym 4,2 V został pokazany na **rysunku 2**. Pierwszy blok pomiarowy jest dłuższy (13 s), a następne trwają tylko 7,7 s i są powtarzane dokładnie co 1 minutę. Pik prądowy przy ponownym włączaniu zasilania

```

LTR-559ALS-01 0x23, PCF85063A 0x51, BME280 0x77, ADS1115 0x48/0x49
Detected devices at I2C1-addresses:
['0x23', '0x48', '0x51', '0x77']
Decimal address: [35, 72, 81, 119]
enviro.startup()
2024-10-14 16:21:57 [info / 126kB] > performing startup
2024-10-14 16:21:57 [debug / 126kB] - running Enviro 0.0.10, MicroPython v1.22.2, enviro v1.22.2 on 2024-03-06
2024-10-14 16:21:58 [info / 119kB] - wake reason: rtc_alarm
2024-10-14 16:21:58 [debug / 117kB] - turn on activity led
2024-10-14 16:21:58 [debug / 113kB] > taking new reading
2024-10-14 16:21:58 [info / 108kB] - seconds since last reading: 53
Rain data: 0.32 %
Rain_volt: 1.34 V
Temp_volt: 1.71 V
Ref_volt: 1.66 V
Wind Direction Voltage: 3.31 V
Wind direction: None
2024-10-14 16:22:00 [info / 106kB] > going to sleep
2024-10-14 16:22:00 [debug / 103kB] - clearing and disabling previous alarm
2024-10-14 16:22:00 [info / 101kB] - setting alarm to wake at 16:23pm
2024-10-14 16:22:00 [info / 99kB] - shutting down
2024-10-14 16:22:00 [debug / 97kB] - on usb power (so can't shutdown). Halt and wait for alarm or user reset instead
2024-10-14 16:22:00 [debug / 95kB] - reset

```

Listing 1. Informacje widoczne po uruchomieniu aplikacji

(zaznaczony na górnym wykresie) to 1,47 A (600 μ s), a w trakcie wykonywania bloku pomiarowego wartość szczytowa to 202 mA (średni prąd: 41 mA). Średni pobór za cały okres pomiarowy 1 minuty (zaznaczony na dolnym wykresie) wynosi tylko 5,37 mA. Pomiędzy blokami pomiarowymi pobór prądu wynosi 30 μ A.

Przy zasilaniu z USB (5V) nie występują szpilki prądowe na początku bloku pomiarowego, średni prąd podczas realizacji całej procedury jest podobny, zaś pomiędzy blokami średni pobór prądu utrzymuje się na poziomie 38 mA.

Podsumowanie

Ocena jakości pracy zestawu elektronicznego wymaga nie tylko sprawdzania jego funkcjonalności, ale również parametrów elektrycznych, zwłaszcza jakości zasilania. Jest to szczególnie ważne, gdy w układzie są obwody cyfrowe zasilane z przetwornicy impulsowej oraz obwody analogowe – często o dużej czułości i wzmocnieniu. W przypadku urządzenia IoT, pracującego w czasie rzeczywistym, dodatkowo bardzo są zależnościami czasowymi.

Analiza niepoprawnego działania pomiaru kierunku wiatru doprowadziła do zastosowania zewnętrznego przetwornika ADC z niskoszumnym zasilaniem rezystorowej drabinki pomiarowej. Rezultatem jest stabilny i dokładny pomiar 8 kierunków geograficznych.

Również układ pomiaru opadów z czujnikiem pojemnościowym, zaadaptowany z pomiaru wilgotności gleby, zmontowany na płytce stykowej i dołączony przy zastosowaniu dosyć długich kabelków, działał zaskakująco stabilnie przy zasilaniu czystym napięciem.

Niezwykle trudno jest przeprowadzić w warunkach domowych wiarygodne próby pomiaru opadów deszczu. Ograniczone testy pokazały, że zmiana pojemności czujnika zależy od rozmiaru powierzchni pokrytej wodą oraz – częściowo – od grubości tej powierzchni. Ustawienie czujnika deszczu pod kątem 30 stopni ułatwia spływanie wody i przyspiesza reakcję czujnika na zmiany natężenia opadów.

Funkcjonalność układu można łatwo rozszerzyć o obliczenia temperatury pojemnościowego czujnika deszczu i sterowanie jego podgrzewaniem. Można również dołączyć inne czujniki, co zostało pokazane w innych artykułach z tej serii.



Rysunek 2. Pobór prądu z baterii, pierwsze trzy bloki pomiarowe

Realizacja oprogramowania bazuje na projekcie „Pomiar parametrów otoczenia, w tym prędkości i kierunku wiatru, opadów deszczu i wilgotności powietrza”, wykonanego w ramach przedmiotu „Systemy wbudowane i oprogramowanie” na Wydziale Elektroniki i Technik Informacyjnych Politechniki Warszawskiej przez zespół w składzie: Weronika Jamróz, Julia Dollani, Kamil Szwagierczak i Mikołaj Pliś.

Henryk A. Kowalski
Instytut Informatyki
Politechnika Warszawska

Bibliografia

- [1] Capacitive Rain Sensor, Radiocontrolli https://tiny.pl/7x6_61xt
- [2] Hacking a Capacitive Soil Moisture Sensor (v1.2) for Frequency Output, April 20, 2021, <https://tiny.pl/d59hm>
- [3] Stacja pogodowa Enviro Weather firmy Pimoroni, EP 04/2024, <https://tiny.pl/d93r1>
- [4] Optymalizacja poboru mocy urządzenia IoT z płytką Raspberry Pi Pico W, EP 05/2024, <https://tiny.pl/d59hg>
- [5] Enviro MicroPython firmware, Pimoroni, <https://tiny.pl/dt49f>
- [6] Profilowanie mocy z zastosowaniem Power Profiler Kit II, EP 05/2022, <https://tiny.pl/d93rd>
- [7] Vaisala RAINCAP® Technology, <https://tiny.pl/6bb8brgq>
- [8] Build Your Own IR Windshield Rain Sensor, Jorge Martinez, Dec. 21, 2017, <https://tiny.pl/vcf2gzf8>
- [9] Ceramic capacitive rain sensor avoids false positives, 2nd November 2016, Telecontrolli, Barney Scott, <https://tiny.pl/9f4z7dq2>
- [10] Moduł z czujnikiem opadów, Kamami, https://tiny.pl/2c_646mb
- [11] Czujnik deszczu – KAMduino UNO oraz Raindrops Module, Mikrokontroler, 04.08.2017, <https://tiny.pl/89-6h9rx>
- [12] Gravity: I2C ADS1115 16-Bit ADC Module, DFR0553, DFRobot, <https://tiny.pl/n6qfj60b>
- [13] ADS1115_mpy, A MicroPython module for the ADS1115 ADC. Wolfgang (Wolle) Ewald, <https://tiny.pl/twv9ztkq9>

Wyświetlacze dotykowe od Artronic

W wielu współczesnych urządzeniach elektronicznych wyświetlacz dotykowy stanowi nie tylko główny, ale wręcz jedyny element interfejsu człowiek-maszyna. Jakość wyświetlanej grafiki – zarówno pod względem odwzorowania kolorów, jak i rozdzielczości, kontrastu czy też kątów widzenia – jest niezwykle ważna dla użytkowników końcowych, ale równie istotne z konstrukcyjnego punktu widzenia okazują się aspekty techniczne, niewidoczne dla odbiorcy, a silnie wpływające na implementację interfejsu GUI. Szczególną pozycję na rynku wyświetlaczy dotykowych zajmują więc moduły ze zintegrowanym sterownikiem graficznym oraz touchpanelem.

Nowy moduł wyświetlacza dotykowego FT811 + FT5446

Do oferty firmy Artronic dołączył ostatnio najnowszy moduł wyświetlacza dotykowego (fotografia 1), składający się z ekranu o przekątnej 7" i rozdzielczości 800×480 pikseli, połączonego z pojemnościowym touchpanelem. Matryca LCD działa pod kontrolą scalonego sterownika graficznego FT811 (fotografia 2) firmy Bridgetek, zaś kontroler dotyku został oparty na układzie FT5446 firmy FocalTech (fotografia 3), w pełni kompatybilnym z FT811. Jest to bardzo korzystna konfiguracja dla programistów pracujących z układami z serii FT811x, pozwala bowiem uniknąć niepotrzebnego zużycia zasobów głównego mikrokontrolera. Komunikacja z modułem odbywa się za pośrednictwem szybkiego i wygodnego w użyciu interfejsu SPI.



Fotografia 1. Nowy moduł wyświetlacza dotykowego na bazie układów FT811 i FT5446

Inteligentne wyświetlacze z interfejsem UART

W wielu przypadkach stopień zaawansowania funkcjonalności urządzenia nie idzie w parze z wymaganą jakością interfejsu graficznego. Innymi słowy, stosowanie rozbudowanego (a więc i kosztownego) mikrokontrolera w przypadku prostszych urządzeń nie ma większego sensu ekonomicznego i technicznego, z drugiej zaś strony mniejszy procesor nie będzie dysponował wystarczającą mocą obliczeniową do płynnej obsługi wysokorozdzielczego ekranu graficznego. W takich sytuacjach doskonale sprawdzają się tzw. wyświetlacze inteligentne, które oprócz niskopoziomowego sterownika graficznego oraz kontrolera dotyku mają także wbudowany procesor obsługujący obszerną pamięć (zwykle w postaci karty microSD) oraz własny system operacyjny.

Więcej informacji:

Artronic sp. j.

81-549 Gdynia, ul. Parkowa 6, tel. 58 668 57 83
biuro@artronic.pl, www.artronic.pl



Fotografia 2. Sterownik graficzny FT811



Fotografia 3. Kontroler dotyku FT5446

Przykładem takiego rozwiązania mogą być ekrany dotykowe marki AV-Display, np. należący do portfolio tego znanego producenta moduł LCD-AG-TFTSD-UART-800480C-70TP (fotografia 4). 7-calowy ekran LCD TFT o rozdzielczości 800×480 px ma wbudowany slot kart microSD i komunikuje się z nadrzędnym mikrokontrolerem za pomocą interfejsu UART. Zestaw prostych w użyciu komend umożliwia sterowanie wysokopoziomowymi funkcjami graficznymi, przykładowo:

`CLS` \n; czyszczenie ekranu

`PIC 0 0 0` \n; obraz nr 0 z karty pamięci w pozycji o współrzędnych (0,0)

`PIC 1 0 0` \n; obraz nr 1 z karty pamięci w pozycji o współrzędnych (0,0)

`STR 360 5 31 TEXT.` \n; wyświetlenie zadanego tekstu w pozycji (360,5) przy użyciu koloru nr 31

Zastosowanie karty microSD w roli pamięci nieulotnej umożliwia zapisanie nawet obszernego zestawu grafik i wyświetlanie ich za pomocą komend przesyłanych w formacie ASCII – bez konieczności implementacji pełnej biblioteki graficznej po stronie procesora nadrzędnego. Moduł wspiera także funkcje animacji, zaawansowane formatowanie tekstu, załączanie i wyłączanie podświetlenia, a także proste w użyciu funkcje graficzne, np. do rysowania punktów, linii, okręgów czy prostokątów. Ten sam interfejs UART jest także używany do zwracania informacji o pozycji naciśnięcia panelu dotykowego – dane są przesyłane w następującym formacie:

`TX Y Xaddr Yaddr` \n, gdzie Xaddr Yaddr to współrzędne „kliknięcia”.



Fotografia 4. Moduł wyświetlacza inteligentnego LCD-AG-TFTSD-UART-800480C-70TP

Sterowanie dotykiem i gestami – przegląd technologii

Sterowanie dotykiem bądź gestami – i to zarówno prostymi machnięciami ręki, jak i rozbudowanymi sekwencjami ruchów obu rąk – przestało już być domeną twórców science fiction, szturmem wdzierając się do naszego codziennego życia. W ślad za rosnącym trendem poprawy ergonomii obsługi urządzeń i systemów elektronicznych idzie także rozwój czujników, modułów i oprogramowania wspierającego detekcję i rozpoznawanie gestów. Niemalą udział ma w tym przypadku sztuczna inteligencja, o której chyba każdy z nas słyszy (w samych tylko środkach masowego przekazu) przynajmniej kilka razy dziennie. W artykule przyglądamy się współczesnym rozwiązaniom przeznaczonym do sterowania dotykiem i gestami w różnych gałęziach techniki.

Kawałek historii, czyli krótko o historii interfejsów dotykowych i HMI opartych na gestach

Sterowanie urządzeniami elektronicznymi przez bezpośrednie wskazywanie konkretnych punktów na ekranie było marzeniem pierwszych użytkowników komputerów. Już w latach 50. XX wieku inżynierowie z MIT opracowali prototyp świetlnego pióra, które zostało chętnie zaadaptowane do coraz popularniejszych systemów marki IBM, Amiga i innych. Jego działanie było bajecznie proste: wskaźnik zbliżony kształtem i rozmiarami do długopisu (fotografia 1) był wyposażony w końcówkę z fotodetektorem, którą użytkownik przykładł do monitora CRT. Podczas przemiatania obrazu wiązka elektronów sukcesywnie rozświetlała luminofor w kolejnych partiach ekranu, a gdy trafiała na czujnik – ten ostatni dawał odpowiedni sygnał z powrotem do komputera. W ten sposób, znając aktualną pozycję plamki światła (X, Y) można było precyzyjnie wskazać położenie świetlnego pióra.



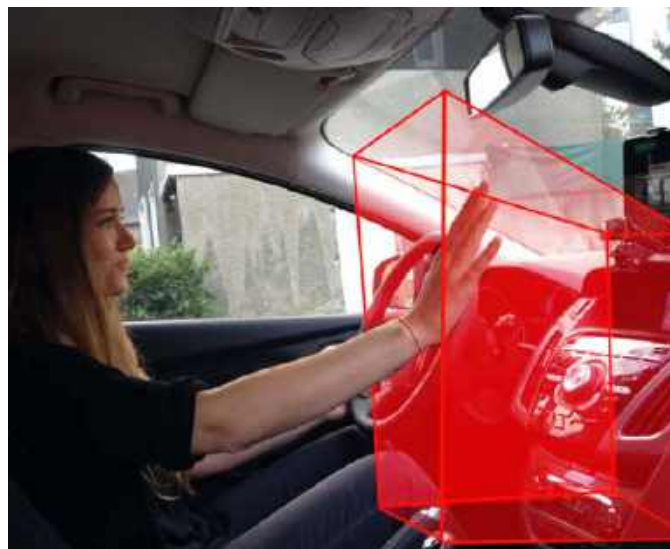
Fotografia 1. Świetlne pióro używane do obsługi komputera IBM – końcówka lat 60. XX wieku (<https://t.ly/PrMli>)

Co ciekawe, już wtedy użytkownicy zaczęli zgłaszać problemy natury ergonomicznej, a nawet zdrowotnej, wynikające z długotrwałego trzymania pióra w wyjątkowo nieergonomicznej pozycji: kończyna górna, zwłaszcza pozbawiona odpowiedniego podparcia, szybko ulegała przemęczeniu, co manifestowało się nawet przypadłością nazywaną mianem „gorilla arm”. Świetlne pióra z czasem praktycznie całkowicie ustąpiły innym, znacznie wygodniejszym w użytku urządzeniom wskazującym: myszom komputerowym i trackballom. Dopiero po latach zaczęły się powoli upowszechniać inne technologie obsługi interfejsów graficznych. I choć ekrany dotykowe są już od około dwóch dekad standardowym rozwiązaniem stosowanym w telefonach komórkowych, tabletach i wielu innych urządzeniach mobilnych (a także ubieralnych), to w przypadku „pecetów” czy laptopów stanowią one mimo wszystko technologię dużo mniej popularną – chyba w głównej mierze właśnie ze względów ergonomicznych. Panele dotykowe królują jednak nie tylko w elektronice konsumenckiej – coraz większy odsetek aparatury pomiarowej (głównie oscyloskopów i analizatorów widma), a także znakomita większość urządzeń medycznych, wyposażonych w różnej wielkości ekrany graficzne, mogą (lub muszą, z braku innych rodzajów interfejsu użytkownika) być sterowane właśnie za pomocą dotyku.

Z czasem popularność zaczęły zyskiwać także rozmaite technologie umożliwiające sterowanie urządzeniami elektronicznymi z pełnym pominięciem fizycznego kontaktu – czyli za pomocą mniej lub bardziej złożonych gestów. Do dziś prawdopodobnie największym rynkiem prostych czujników zbliżeniowych pozostają dyspensery mydła i środków do dezynfekcji rąk czy też bezdotykowe baterie umywalkowe, ale coraz głośniej mówi się o nadchodzącej rewolucji



Fotografia 2. Bateria umywalkowa z bezdotykowym czujnikiem (<https://t.ly/CfZRb>)



Fotografia 3. Przykład sterowania osprzętem pokładowym pojazdu za pomocą gestów (<https://t.ly/R37am>)

w obsłudze instalacji pokładowych w samochodach osobowych (fotografia 3). A teraz ciekawostka: prognozy biznesowe wskazują, że globalny rynek technologii rozpoznawania gestów w aplikacjach motoryzacyjnych już w 2020 roku był wyceniany na niecały miliard dolarów, a w roku 2030 będzie dobiegał do 4,5 miliarda! Nie dziwi zatem, że przeglądając strony internetowe producentów półprzewodników co rusz natrafiamy na infografiki kuszące nas wymyślnymi systemami sterowania wyposażeniem pokładowym za pomocą ruchów ręki.

Ekspresowy przegląd technologii wyświetlaczy dotykowych

Przez kilka dekad rozwoju ekranów graficznych – równolegle z dopracowywaniem istniejących i poszukiwaniem nowych metod wyświetlania obrazu – powstawał także szereg technologii pozwalających na precyzyjną i szybką detekcję dotyku. Pokróćce omówimy najważniejsze z nich, zwracając uwagę na aspekty implementacyjne oraz zalety i wady, rozpatrywane z użytkowego punktu widzenia.

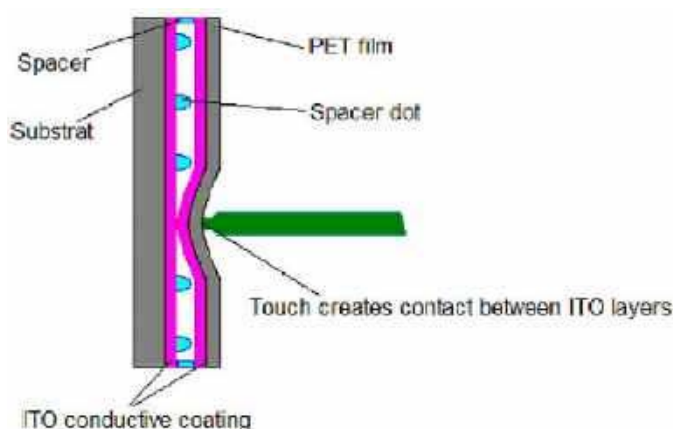
Touchpanele rezystancyjne

Technologia rezystancyjna opiera się na pomiarze zmian oporności pomiędzy elektrodami, połączonymi galwanicznie z dwiema warstwami przewodzącymi. Te ostatnie są przezroczyste (wykonane z tlenku indy i cyny, ITO) i umieszczone w niewielkiej odległości od siebie, pozostają „domyślnie” rozdzielone przez warstwę elementów dystansowych w postaci niewidocznych gołym okiem „kropek”. Naciśnięcie warstwy elastycznej powoduje jej ugięcie i zwarcie z powłoką ITO umieszczoną na sztywnym podłożu panelu (rysunek 1).

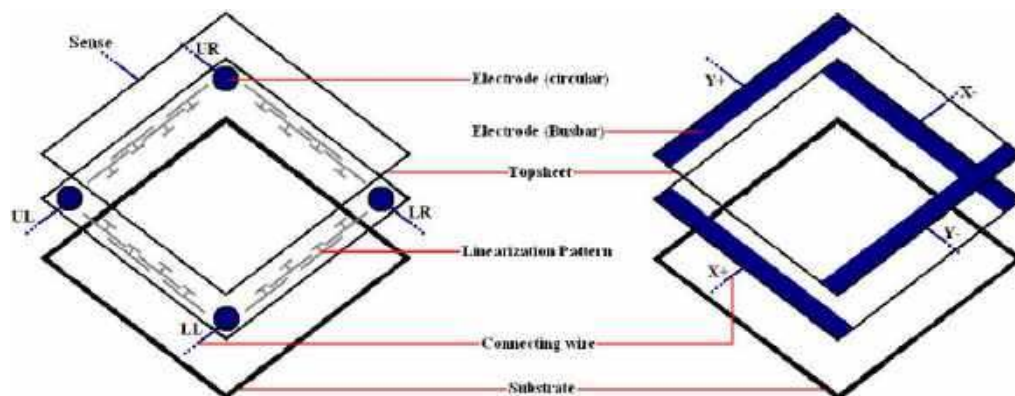
Rezystancja może być zmierzona na kilka sposobów, a najczęściej spotykane są połączenia 4- oraz 5-przewodowe (rysunek 2b). W konfiguracji 4-przewodowej elektrody paskowe umieszcza się na wszystkich czterech krawędziach ekranu, przy czym elementy poziome są połączone elektrycznie z jedną z warstw ITO, zaś pionowe – z drugą. W procesie odczytu najpierw zasilą się pierwszą warstwę poprzez przyłożenie napięcia pomiędzy należące do niej elektrody, a następnie odczytuje się napięcie w punkcie dotyku, używając pozostałych (niezasilonych) elektrod niczym... suwaka potencjometru. W wyniku pomiaru otrzymujemy wartość napięcia, która – po uwzględnieniu współczynników kalibracyjnych – pozwala na dokładne określenie pozycji w osi prostopadłej do zasilonych elektrod. Następnie procedura jest wykonywana analogicznie na drugiej warstwie, tj. po funkcjonalnej zamianie zestawów elektrod. W ten sposób można określić położenie punktu styku obu warstw ITO – do realizacji pomiaru wystarczy zatem... odpowiedni multiplexer oraz przetwornik ADC. Nic nie stoi oczywiście na przeszkodzie, by uprościć sobie pracę i zastosować scalony kontroler – przykład takiego rozwiązania można zobaczyć na rysunku 3.

Nieco inaczej wygląda sprawa w przypadku paneli 5-przewodowych (rysunek 2a). Tutaj odczyt jest dokonywany

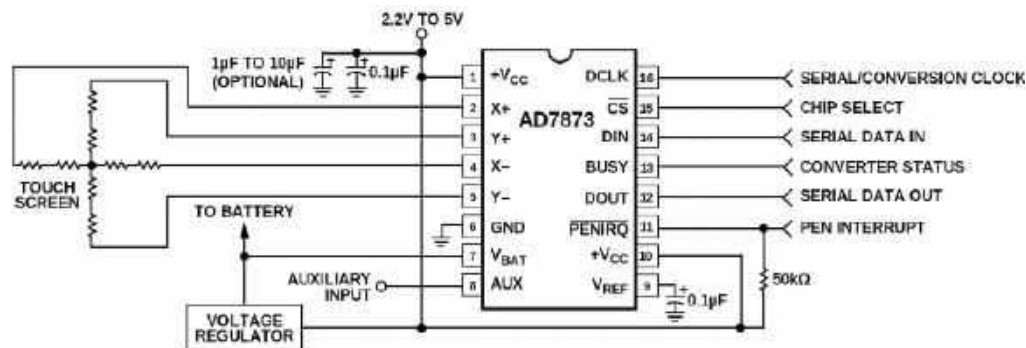
po zasileniu par sąsiadujących elektrod: w celu odczytu w osi X napięcie jest przykładane między elektrody lewe (np. potencjał masy), a prawe (np. potencjał zasilania), a w drugiej fazie – pomiędzy elektrody górne, a dolne. Odczyt jest dokonywany z piątej elektrody oznaczonej jako „Sense” na rysunku 2a i podłączonej do górnej warstwy ITO. Co ciekawe, aby umożliwić dokładne określenie położenia punktu kontaktowego, konieczne jest dodanie specjalnych ścieżek pomiędzy elektrody leżące przy danej krawędzi wyświetlacza, a których celem jest linearyzacja gradientu napięcia na obszarze całego touchpanelu. Z tego też względu ekrany dotykowe z 5-przewodowym panelem dotykowym są trudniejsze do wyprodukowania, ale oferują istotną zaletę: okazują się bowiem mniej wrażliwe na uszkodzenia struktury ITO, co znakomicie podnosi ich niezawodność względem wersji 4-przewodowej. Technologia rezystancyjna (ogólnie) jest ponadto odporna na zanieczyszczenia stałe i wodę (krople znajdujące się na powierzchni panelu), umożliwia także obsługę za pomocą rękawiczek lub dowolnych przedmiotów pozbawionych ostrych krawędzi (np. rysika lub tępej końcówki długopisu). Do wad



Rysunek 1. Ilustracja zasady działania rezystancyjnego panelu dotykowego (https://t.ly/c_wBi)



Rysunek 2. Porównanie układów elektrod w panelach rezystancyjnych. a) 5-elektrowy, b) 4-elektrowy (https://t.ly/c_wBi)



Rysunek 3. Przykładowa implementacja kontrolera 4-przewodowego, rezystancyjnego panelu dotykowego przy użyciu układu AD7873 (<https://t.ly/pqtLw>)



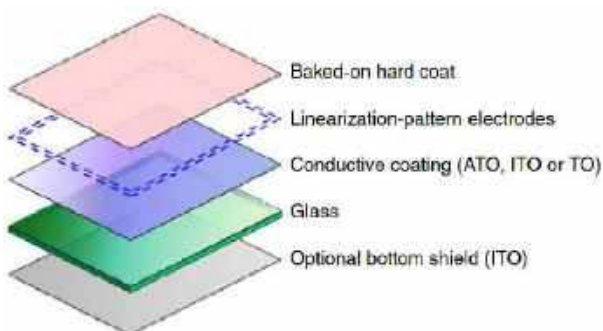
Fotografia 4. Przykładowe moduły kontrolerów wielodotykowych paneli rezystancyjnych (<https://t.ly/-iA-r>)

należy natomiast zaliczyć mniejszą czułość na dotyk i nieco gorszą widoczność w silnym świetle słonecznym (z uwagi na złożoną strukturę touchpanelu, tłumiącą w widocznym stopniu światło emitowane przez wyświetlacz).

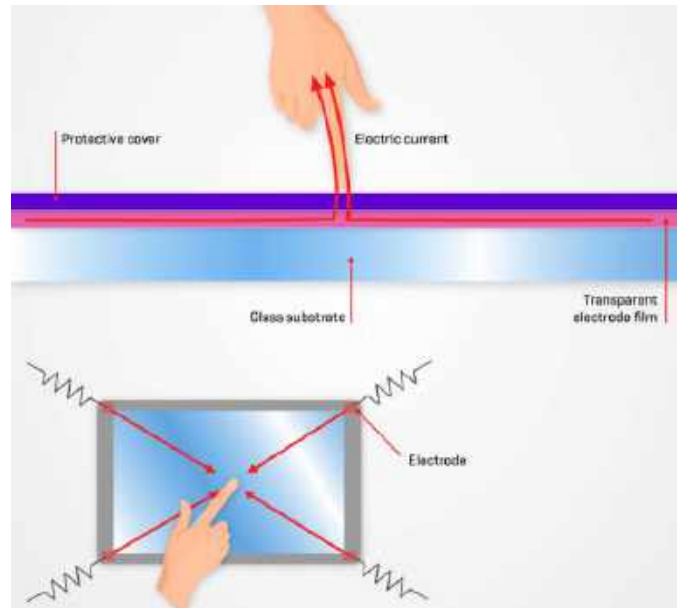
Patrząc na schematy pokazane na rysunku 2 nietrudno dojść do słusznego wniosku, że klasyczna technologia rezystancyjna uniemożliwia realizację funkcji multi-touch. Tak jest w istocie – dotknięcie ekranu w dwóch lub kilku punktach jednocześnie zaburzy pomiar i spowoduje wygenerowanie fałszywych, „pośrednich” wyników. Warto jednak dodać, że istnieją sposoby obejścia tego problemu – wymagają one jednak podziału warstw ITO na mniejsze pola, co niebawem komplikuje przemiatanie ekranu i wymusza znaczną rozbudowę sprzętową kontrolera. Widać to doskonale na **fotografii 4**, pokazującej przykładowe sterowniki panelu multi-touch – imponująco duże złącza krawędziowe to jedyny sposób na podłączenie obszernego zestawu elektrod, biegnących od touchpanelu w zintegrowanej z nim taśmie FPC.

Panele pojemnościowe

Wśród dotykowych paneli pojemnościowych także można wyróżnić dwie główne technologie, tym razem jednak różniące się od siebie w znacznie większym stopniu, niż miało to miejsce w opisanych wcześniej konstrukcjach rezystancyjnych. Powierzchniowe panele pojemnościowe, w skrócie: SCT (ang. Surface Capacitive Touch) mają konstrukcję zdecydowanie najprostszą w swojej klasie, stąd ich najpowszechniejsze zastosowania to większe ekrany dotykowe o ograniczonych wymagach w zakresie rozdzielczości i dokładności lokalizacji punktów dotykowych. Ich konstrukcja opiera się na pojedynczej, jednolitej warstwie przewodzącej, wyposażonej w ścieżki linearyzujące na bazie srebra i podłączonej do czterech elektrod, znajdujących się w rogach ekranu (**rysunek 4**). Przyłożenie napięcia przemiennego do narożnych elektrod powoduje



Rysunek 4. Budowa powierzchniowego panelu pojemnościowego (SCT). Źródło: <https://t.ly/u2BiQ>



Rysunek 5. Zasada działania powierzchniowego panelu pojemnościowego (<https://t.ly/FoJ06>)

wytworzenie pola elektrostatycznego. Przewody doprowadzone do poszczególnych punktów nie są jednak ze sobą zwarte: kontroler wytwarza cztery przebiegi o identycznej fazie, częstotliwości i amplitudzie. Przyłożenie palca lub rysika przewodzącego (trzymanego przez użytkownika) powoduje zaburzenie równowagi pola, przez co do poszczególnych elektrod dopływa prąd o nieco różniących się natężeniach – amplituda poszczególnych prądów jest bowiem proporcjonalna do odległości punktu kontaktowego od danej

REKLAMA

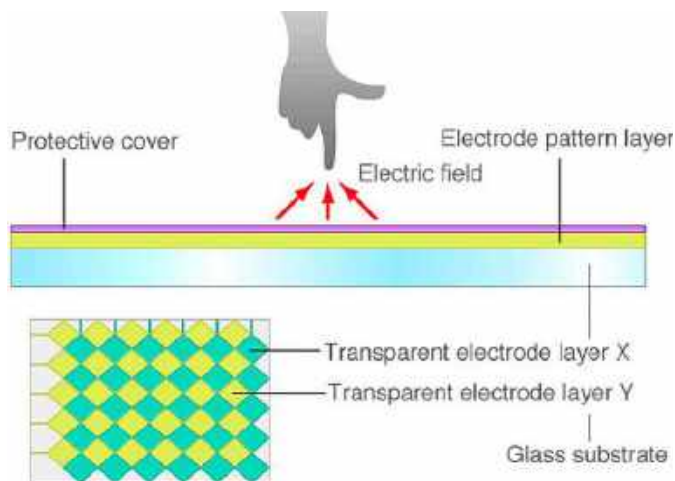
Moduły rozszerzeń do mikrokontrolerów, w tym do Raspberry Pi Pico



Wyświetlacze HMI firm Nextion, DWIN



Elty sp. z o.o.
43-518 Zabrzeg, ul. Miliardowicka 17 A
elty@elty.pl, <https://elty.pl/>



Rysunek 6. Budowa projekcyjnego panelu pojemnościowego (<https://t.ly/WoyDH>)

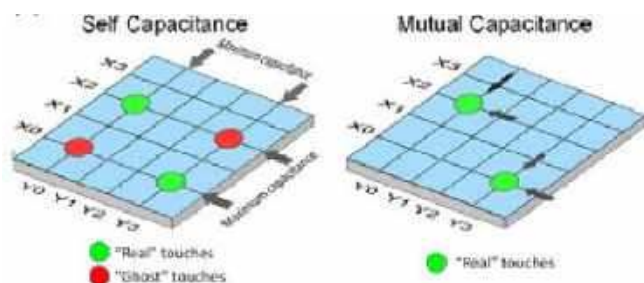
elektrody (rysunek 5). Jak nietrudno się domyślić, z uwagi na zastosowanie pojedynczej warstwy i tylko czterech elektrod pracujących niejako „współ w zespół”, nie da się zrealizować w ten sposób panelu zdolnego do rozróżnienia wielu punktów dotykowych jednocześnie. Innym problemem jest wymóg relatywnie stabilnej masy, stanowiącej odniesienie dla układu pojemności: ekran-użytkownik-ziemia. Z tego też względu ekrany w technologii SCT lepiej nadają się do instalacji stacjonarnych (np. infokiosków), niż do urządzeń przenośnych.

Diametralnie inne możliwości daje – znacznie bardziej rozbudowana – technologia pojemnościowych paneli projekcyjnych (rysunek 6). W tym przypadku stosowane są dwie warstwy elektrod, ułożonych w układzie macierzowym (X-Y), przemiatanym kolejno przez sterownik.

Gwoli ścisłości należy dodać, że istnieją dwie główne odmiany techniki projekcyjnej. Pierwsza z nich opiera się na pomiarze pojemności własnej danej elektrody, czyli – innymi słowy – odczyt jest dokonywany w odniesieniu do masy urządzenia. Przyłożenie palca do ekranu w danym miejscu zwiększa widzianą przez kontroler dotyku pojemność, zatem położenie punktu kontaktowego jest wyznaczane na drodze poszukiwania maksimum pojemności (a ściślej rzecz ujmując: największej zmiany względem pojemności bazowej).

Druga metoda polega na pomiarze pojemności pomiędzy dwiema sąsiadującymi elektrodami. W tym przypadku dotknięcie ekranu powoduje „zabranie” części ładunku przez pojemność reprezentowaną przez ciało użytkownika – kontroler musi zatem poszukiwać pary elektrod o pojemności niższej niż pojemność bazowa.

Niezależnie od zastosowanej odmiany technologii, projekcyjne panele pojemnościowe należą do najczulszych i najbardziej responsywnych. Mało tego – ze względu na zastosowany układ macierzowy elektrod umożliwiają one realizację funkcji multi-touch, choć taka implementacja nie zawsze jest trywialna. Kontroler musi



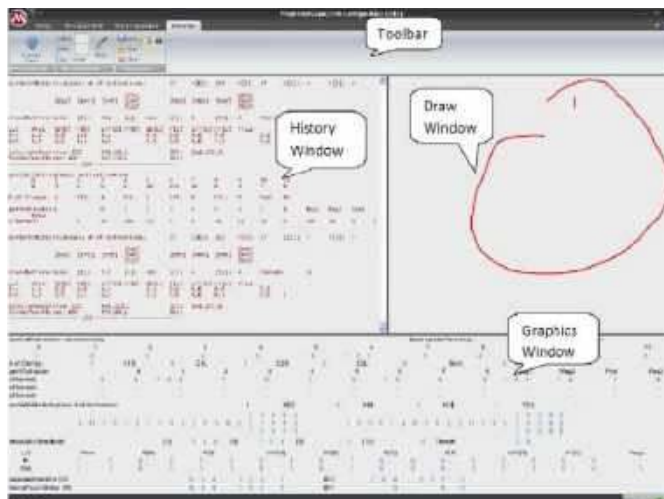
Rysunek 7. Wyjaśnienie procesu powstawania artefaktów podczas obsługi funkcji multi-touch w projekcyjnych panelach pojemnościowych z technologią odczytu pojemności własnej elektrod (<https://t.ly/u2BiQ>)

bowiem uwzględnić zjawisko „duchów”, czyli artefaktów spowodowanych „rzutowaniem” punktów dotyku w dwóch ortogonalnych kierunkach (rysunek 7) – problem dotyczy na szczęście tylko macryc działających w oparciu o technologię pomiaru pojemności własnej, gdyż w przypadku paneli mierzących pojemność wzajemną takie zjawisko nie występuje.

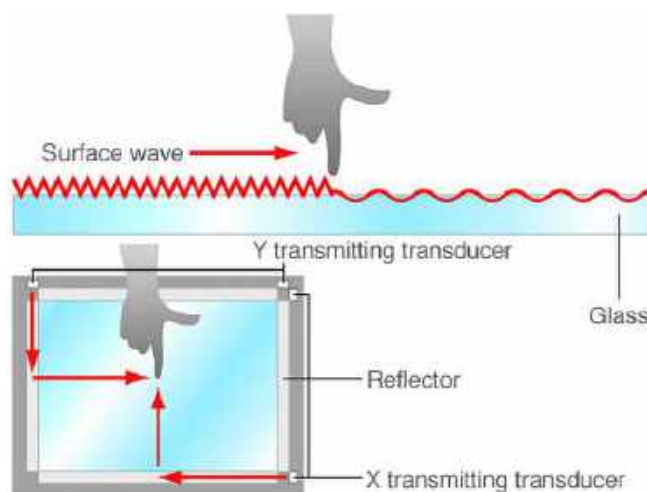
Warto w tym miejscu dodać, że technologia projekcyjna jest stosowana z powodzeniem nie tylko w budowie paneli dotykowych współpracujących z ekranami graficznymi. Z równym powodzeniem można ją również zastosować do budowy touchpadów – w tej realizacji jest ona zresztą diametralnie prostsza, cały układ może bowiem bazować na odpowiednio zaprojektowanej płytce drukowanej. Doskonałym przykładem takiej implementacji jest zestaw ewaluacyjny marki Microchip o nazwie DM160219 (fotografia 5). Ten sam producent przygotował także zestaw o numerze DM160211, przeznaczony do eksperymentów z kontrolerem dotyku opartym



Fotografia 5. Zestaw ewaluacyjny DM160219 firmy Microchip (<https://t.ly/2xDGU>)



Rysunek 8. Widok okna aplikacji współpracującej z zestawem DM160211 (<https://t.ly/nYn5b>)



Rysunek 9. Zasada działania ekranu dotykowego wykonanego w technologii SAW (<https://t.ly/YjU3U>)



Fotografia 6. Zestaw ewaluacyjny DM160211 firmy Microchip
(<https://t.ly/irwvL>)

na mikrokontrolerze PIC16F707 i pozwalający nie tylko na efektywne rysowanie dowolnych kształtów w czasie rzeczywistym (w specjalnie przygotowanej aplikacji demonstracyjnej), ale także na odczyt i wizualizację (w różnych trybach) parametrów odczytywanych z procesora znajdującego się na płycie (**rysunek 8**).

Zdolność do rozpoznawania dotyku wielopunktowego (a więc także wsparcie obsługi gestów) oraz doskonała szybkość, responsywność i precyzja działania paneli projekcyjnych spowodowała, że technologia ta bezsprzecznie zdominowała najbardziej wymagające (pod względem doświadczeń użytkownika) obszary aplikacji elektroniki konsumenckiej – głównie rynek urządzeń mobilnych, w tym smartfonów i tabletek.

Akustyczna technologia powierzchniowa

Panele dotykowe można budować także w oparciu o powierzchniową falę akustyczną (ang. Surface Acoustic Wave, SAW), tzw. falę Rayleigha – propagującą praktycznie tylko na powierzchni oraz na niewielkim obszarze w głąb materiału. W praktyce stosowane są fale ultradźwiękowe o częstotliwości rzędu 5 MHz – stosunkowo łatwe do wygenerowania i odebrania za pomocą przetworników piezoelektrycznych. Na krawędziach ekranu montowane są podłużne reflektory akustyczne, zaś w narożnikach – dwie pary przetworników (po jednym nadawczym i odbiorczym, osobno dla osi X oraz Y – patrz **rysunek 9**). Fale propagują – za sprawą umieszczonych pod kątem 45° reflektorów – najpierw w osi równoległej do kierunku drgań, a następnie prostopadłe do niej, przez całą powierzchnię panelu, aż do przeciwnego zestawu reflektorów. Tam ulegają odbiciu w stronę przetwornika odbiorczego, przy czym – co jest niezwykle ważne dla działania tej technologii – fale odbite od początku reflektora rozpoczynają i kończą swój bieg wcześniej, niż te odbite od jego końca. W ten sposób odległość w danej osi jest mapowana w domenie czasu – a zatem dotknięcie panelu palcem lub miękkim przedmiotem, powodujące częściowe, lokalne stłumienie amplitudy drgań, będzie widoczne w sygnale płynącym z przetwornika odbiorczego jako chwilowy „zapad” w amplitudzie napięcia.

Panele zbudowane na bazie technologii SAW mają szereg zalet – zapewniają doskonałą przezierność (jedynym elementem, który znajduje się przed właściwym wyświetlaczem, jest po prostu... szyba), odporne na zarysowania i uderzenia (zależnie tylko od wytrzymałości szkła), relatywnie proste w konstrukcji, niezawodne i skalowalne, przez co mogą być stosowane zarówno w niewielkich,

kilkucalowych ekranach, jak i sporych wyświetlaczach o przekątnej na poziomie kilkudziesięciu cali. Wysoka odporność środowiskowa i mechaniczna pozwala na stosowanie dotykowych paneli akustycznych w aplikacjach outdoorowych (**fotografia 7**).

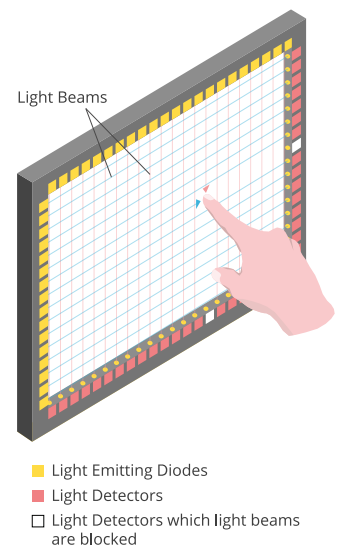
Panele dotykowe na podczerwień

Ostatnia z technologii dotykowych, których nie mogło zabraknąć w niniejszym artykule, jest zdecydowanie najprostsza do zrozumienia, a zarazem... najbardziej złożona pod względem konstrukcyjnym. Mowa o panelach opartych na dwuwymiarowych kurtynach wiązek podczerwieni (**rysunek 10**). Zasada działania jest banalnie prosta – zbliżenie palca na odległość wystarczającą do przerywania wiązek światła powoduje, że określone fotodetektory, zamontowane na krawędziach ramki (**fotografia 8**), są przesłaniane. Przerwanie danej bariery optycznej jest jednoznaczny wskazaniem współrzędnej w osi X lub Y.

Opisywana technologia z powodzeniem obsługuje funkcję multi-touch, a ponadto idealnie nadaje się do realizacji ogromnych paneli o przekątnych na poziomie nawet 150 cali (!). Dlatego też duże infokioski (**fotografia 9**), interaktywne instalacje rozrywkowe czy edukacyjne oraz wszelkiego rodzaju inne rozwiązania wymagające wykrywania dotyku na sporej powierzchni są oparte właśnie na technologii kurtyn IR. Mało tego – optoelektroniczna rama może być z niebywałą łatwością



Fotografia 7. Przemysłowy ekran dotykowy z panelem wykonanym w technologii SAW
(<https://t.ly/0l1c2>)



Rysunek 10. Zasada działania optycznego panelu dotykowego
(<https://t.ly/foqGe>)



Fotografia 8. Ramka z rzędami zamontowanych naprzemiennie fotoelementów i diod LED IR
(<https://t.ly/iRjgP>)



Fotografia 9. Przykładowa aplikacja podczerwonego panelu dotykowego
(<https://t.ly/7h73k>)



Fotografia 10. Montaż kurtyny IR do istniejącego ekranu
(<https://t.ly/JWx3z>)

zamontowana do dowolnego ekranu wielkoformatowego, przez co producenci i integratorzy systemów zyskują ogromne pole do popisu pod względem wyboru dostawcy i modelu wyświetlacza (**fotografia 10**).

A jakie są wady paneli dotykowych IR? Podstawowa i zarazem największa wada wynika wprost z samej technologii: uzyskanie relatywnie wysokiej rozdzielczości wymaga stosowania setek, a nawet tysięcy fotoelementów i diod LED, montowanych na długich i dość rozbudowanych płytkach drukowanych. To wprost idealny przepis na drogą technologię, która jest równie skalowalna, co wrażliwa na ceny i liczbę zastosowanych komponentów.

Systemy detekcji gestów

Zróżnicowanie stosowanych w praktyce technologii wykrywania gestów jest jeszcze większe, niż w przypadku opisanych wcześniej paneli dotykowych. Scharakteryzujemy zatem – w telegraficznym skrócie – dostępne opcje.

Techniki optyczne

Rozwiązania optoelektroniczne zdecydowanie królują wśród wszystkich znanych technik detekcji i rozpoznawania gestów. Najprostszą realizację można wykonać z użyciem zwykłego transoptora odbiciowego i odpowiedniego oprogramowania – w ten sposób można z powodzeniem wykrywać nie tylko sam fakt zbliżenia/oddalenia dłoni do/od czujnika, ale także (choć bardziej jakościowo, niż ilościowo) określać jej odległość od sensora.



Fotografia 11. Moduł czujnika gestów na bazie nadajników IR i scalonego odbiornika podczerwieni (<https://t.ly/HyFW2>)

Praca z zasilaniem emitera (diody nadawczej IR) prądem stałym zdecydowanie nie byłaby najlepszą opcją, a to ze względu na znaczną podatność na zmienność warunków oświetlenia zewnętrznego, nawet przy zastosowaniu skutecznego filtra optycznego na pasmo podczerwieni. Zdecydowanie lepiej zastosować światło modulowane – a skoro tak, to (prawie) nic nie stoi na przeszkodzie, by do wykrywania światła odbitego od dłoni zastosować scalony odbiornik podczerwieni, spotykany powszechnie w systemach zdalnego sterowania IR. Jeśli wzbogacimy go o dwa, załączane naprzemiennie emitery, uzyskamy możliwość wykrywania przybliżonej odległości, ale także... kierunku przesuwania dłoni w jednej osi. Dokładnie tak działa jeden z modułów marki SparkFun – ZX Gesture Sensor (**fotografia 11**).

Na przestrzeni ostatniej dekady znacznie większą popularność zdobyły scalone czujniki zbliżeniowe, wyposażane często w funkcje pomiaru natężenia i koloru światła otoczenia. Przykładem mogą być miniaturowe sensory APDS-9960 marki Avago (**fotografia 12**) – ponieważ jednak na łamach EP niejednokrotnie prezentowaliśmy już hybrydowe rozwiązania czujników ALS, sensorów zbliżeniowych i detektorów koloru oraz gestów, zainteresowanych Czytelników zachęcamy do zapoznania się ze stosownymi materiałami archiwalnymi, znajdującymi się na portalu <https://ep.com.pl>.

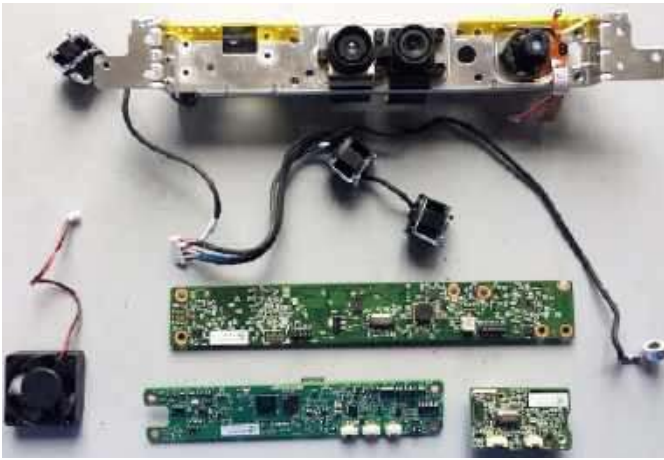
Upowszechnienie kompaktowych dalmierzy o konstrukcji opartej na laserach VCSEL oraz detektorach SPAD spowodowało znaczny wzrost udziału „prawdziwych” czujników odległości (zdolnych do określenia rzeczywistej odległości obiektu wyrażonej w milimetrach) w systemach detekcji gestów. O ile jednak najprostsze układy dToF dobrze nadają się tylko do rozpoznawania ruchów wykonywanych w osi optycznej czujnika (bliżej/dalej), to już matrycowe sensory dają znacznie większe możliwości. Szczególne zasługi w upowszechnianiu tej technologii optoelektronicznej ma firma ST Microelectronics, której najnowszy czujnik – VL53L8CX (**fotografia 13**) – oferuje aż 64-polowy detektor SPAD (matryca 8×8), umożliwiający rozpoznawanie szeregu złożonych sekwencji ruchu dłoni, we wszystkich trzech płaszczyznach.



Fotografia 13. 64-polowy sensor dToF marki ST Microelectronics (<https://t.ly/2SjtH>)



Fotografia 14. Przykład komercyjnego modułu kamery iToF (<https://t.ly/NsYtA>)

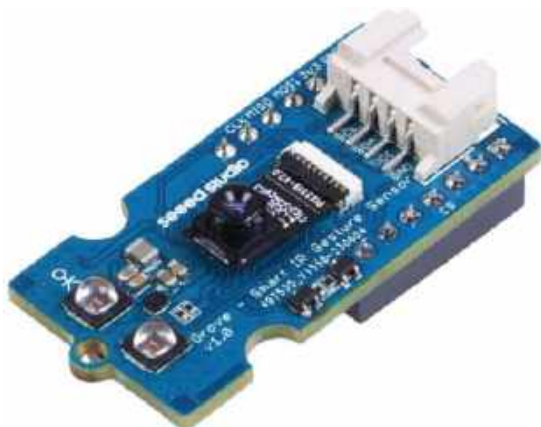


Fotografia 15. Wnętrze czujnika Microsoft Xbox 360 Kinect (<https://t.ly/fGjh5>)

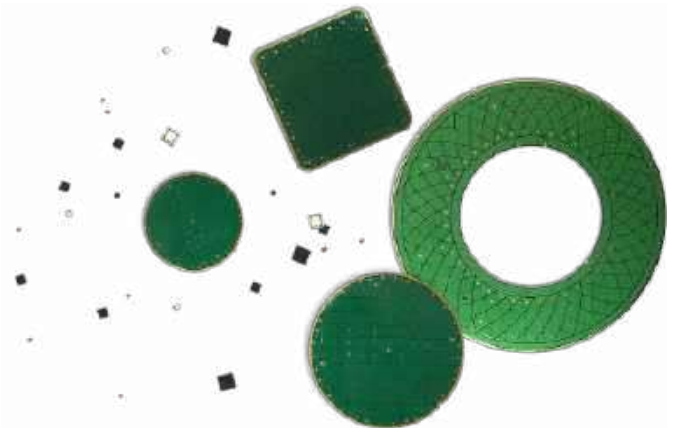
Coraz więcej doniesień prasowych dotyczy także stosowania jeszcze bardziej złożonych detektorów głębi, czyli kamer iToF. Technologia ta – jakkolwiek najbardziej obiecująca pod względem możliwości obrazowania trójwymiarowego – okazuje się jednak wciąż dość droga i relatywnie mało popularna. Należy jednak spodziewać się stopniowego spadku cen i wzrostu rozdzielczości, co – w połączeniu z ekspansją metod sztucznej inteligencji i przetwarzania brzegowego – z pewnością doprowadzi do znacznego rozszerzenia praktycznych aplikacji kamer iToF. Zwłaszcza, że już teraz dostępne są relatywnie kompaktowe moduły o coraz większych możliwościach (fotografia 14).

Co ciekawe, na długo przed wprowadzeniem na rynek pierwszych kamer 3D dostępne były (i to w przystępnych, nawet dla użytkowników prywatnych, cenach detalicznych) inne systemy obrazowania trójwymiarowego. Do sztandarowych przykładów należy sensor Kinect, który na pewien czas zrewolucjonizował nie tylko sposób myślenia o grach komputerowych, ale także trafił do licznych projektów badawczych w zakresie informatyki (wizyjne śledzenie ruchu), medycyny (komputerowe wspomaganie chirurgii) czy przemysłu. W tym przypadku jednak nie mieliśmy już do czynienia z „prostym”, pojedynczym czujnikiem scalonym, a ze złożoną platformą zawierającą dwie kamery oraz projektor światła strukturalnego, współpracujący z detektorem obrazu IR (fotografia 15).

W tym miejscu warto dodać, że producenci modułów OEM także nie próżnują i maksymalnie eksploatują dostępne technologie – nawet te o znacznie dłuższej historii. Przykładem mogą być miniaturowe kamery marki PixelArt, współpracujące z wyspecjalizowanymi procesorami obrazu i rozpoznające gesty na podstawie dwuwymiarowego obrazu IR (fotografia 16).



Fotografia 16. Przykładowy moduł przeznaczony do rozpoznawania gestów na drodze obróbki obrazu z miniaturowej kamery IR (<https://t.ly/uE81c>)



Fotografia 17. Pojemnościowe interfejsy dotykowe oparte na PCB oraz przeznaczone do ich obsługi kontrolery scalone z oferty firmy Azoteq (<https://t.ly/e4GNY>)

Techniki pojemnościowe

Dopracowanie technik niezawodnej detekcji dotyku za pomocą układów elektrod wpłynęło na wzrost popularności interfejsów HMI opartych na przyciskach, pokrętkach czy suwakach (fotografia 17), realizowanych bez ani jednego elementu mechanicznego. Wystarczy przymocować odpowiednio zaprojektowaną płytkę drukowaną do wewnętrznej powierzchni obudowy – metoda prosta, skuteczna, niezwykle estetyczna i tania w implementacji. Interfejsy dotykowe, pracujące bez powiązania z jakimkolwiek wyświetlaczem, zostały dopracowane do granic technicznych możliwości i na stałe zadomowiły się rynku elektroniki konsumenckiej, medycznej, w sprzęcie RTV i AGD oraz wielu innych obszarach technologii.

Rozwinięcie znanych wcześniej sposobów detekcji dotyku – głównie w kierunku dalszego zwiększania czułości, ulepszania metod redukcji zakłóceń i artefaktów oraz odpowiedniego przetwarzania danych płynących z przetworników pojemności – doprowadziło do opracowania rozwiązań umożliwiających bezdotykowe sterowanie urządzeniami elektronicznymi. Rozwiązania te opierają się na detekcji zaburzeń pola elektrycznego, które generowane jest w sposób zbliżony do tego, który zaprezentowaliśmy przy okazji opisu powierzchniowych paneli dotykowych (rysunki 11a i 11b). Firma Microchip wprowadziła do swojej oferty zaawansowane, wielokanałowe układy do rozpoznawania gestów, bazujące na autorskiej technologii GestIC. W zależności od rodzaju implementacji i warunków eksploatacyjnych, zasięg detekcji rozciąga się od kilkunastu do nawet 200 mm, czyli... dogania pod tym względem rozwiązania oparte na czujnikach optoelektronicznych. Co ważne – wciąż bez konieczności wykonywania jakichkolwiek dodatkowych otworów i okien optycznych w obudowie. Przykładowo

REKLAMA

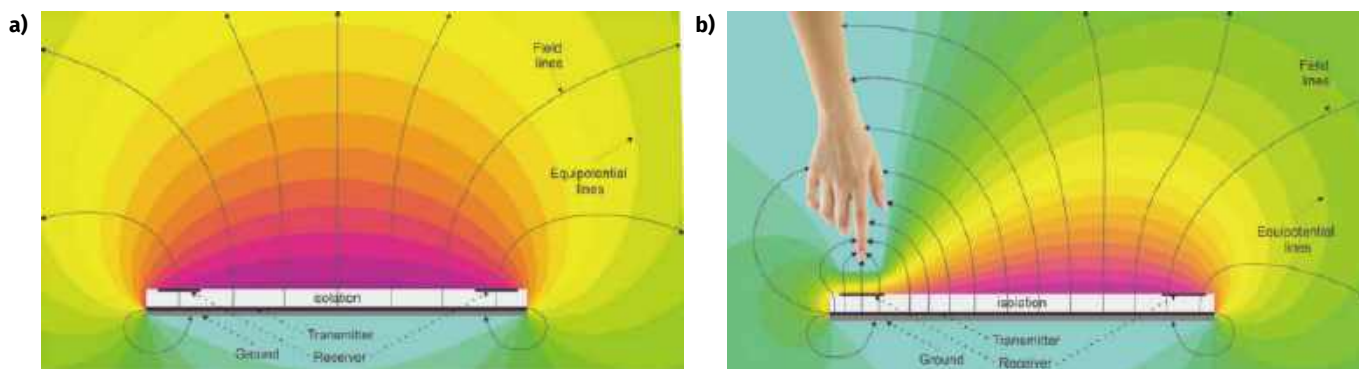


to miejsce, które
łącząc doświadczenie
z innowacyjnością sprawia, że Twoje
pomysły nabierają życia.

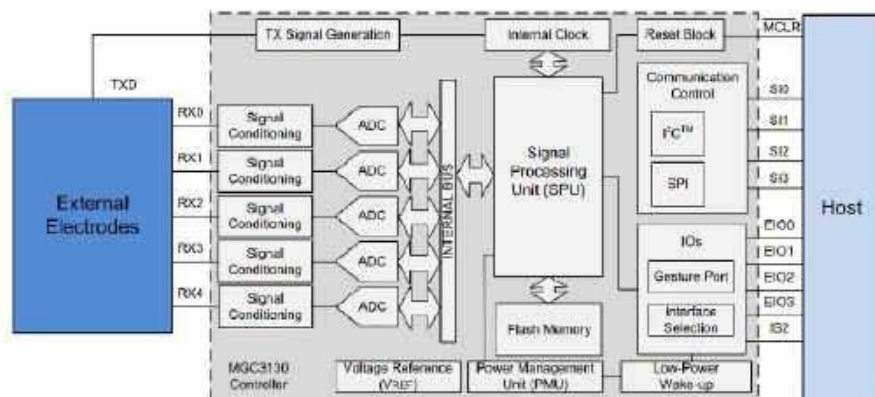


✉ bornico@bornico.com.pl 🌐 www.bornico.com.pl

☎ +48 517 312 709 | +48 517 312 419



Rysunek 11. Układ linii sił pola elektrycznego wokół pojemnościowego czujnika gestów: a) pole niezaburzone, b) pole zaburzone przez zbliżoną dłoń użytkownika (https://t.ly/g00_D)



Rysunek 12. Schemat blokowy układu MGC3130 firmy Microchip (<https://t.ly/EUTir>)

– czułość układu MGC3130 (rysunek 12) dochodzi do 1 fF, szybkość odczytów – do 200 pozycji/sekundę, zaś rozdzielczość przestrzenna – aż do zawrotnej wartości 150 dpi (!). A to wciąż nie wszystko – układ MXG3141 powstał z myślą o rozbudowie istniejących systemów, opartych na ekranach pojemnościowych, o funkcję rozpoznawania trójwymiarowych gestów. Elektrody touchpanelu są współdzielone przez dwa układy: wspomniany MXG3141 oraz właściwy kontroler dotyku 2D (rysunek 13).

Techniki mikrofalowe

W ostatnich latach da się zauważyć znaczący rozwój technik mikrofalowych w paśmie 60 GHz. Miniaturowe radary FMCW (z modulowaną częstotliwościowo falą ciągłą) stały się niezwykle intensywnie eksploatowanym obszarem technologii, zwłaszcza w branży motoryzacyjnej oraz automatycznej budowlanej. Możliwość bezdotykowego pomiaru nawet niewielkich przemieszczeń umożliwia sprawną detekcję gestów, co więcej – podobnie, jak w przypadku systemów pojemnościowych – bez konieczności wprowadzania jakichkolwiek elementów poza obudowę urządzenia.

Czujniki radarowe mają jednak także dodatkowe przewagi nad sensorami pojemnościowymi – nie są bowiem uzależnione od sprzężenia z ziemią, przez co detekcja może odbywać się z równym powodzeniem w dowolnej implementacji, a ponadto okazują się nieporównanie mniejsze – przykładowy, zintegrowany czujnik radarowy 60 GHz marki Infineon (model BGT60TR13C

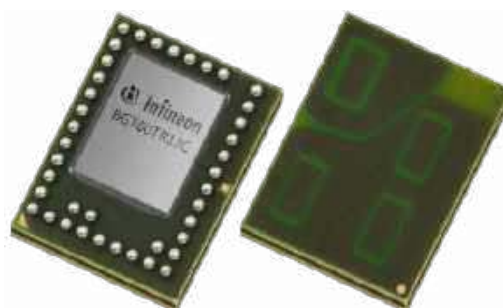
– fotografia 18) ma wymiary zaledwie 6,5×5,0×0,9 mm², a w odróżnieniu od układów bazujących na polu elektrycznym, nie wymaga żadnych dodatkowych elementów peryferyjnych. Anteny są już bowiem wbudowane w sam układ – w tym przypadku są to trzy anteny odbiorcze (ustawione na kształt litery „L”, co pozwala na pomiar ruchów w dwóch wzajemnie ortogonalnych osiach) oraz jedną nadawczą.

Podsumowanie

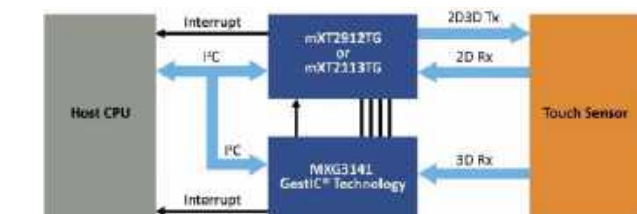
W artykule dokonaliśmy ekspresowego przeglądu dostępnych na rynku technologii dotykowego oraz bezdotykowego sterowania urządzeniami elektronicznymi. Świadomie pominęliśmy przy tym szereg bardziej niszowych rozwiązań – gwoździ ścisłości należy bowiem dodać, że w literaturze naukowej oraz na stronach producentów można natrafić na rozliczne przykłady mało popularnych technik, np. w zakresie ultradźwiękowej detekcji gestów czy też indukcyjnych klawiatur dotykowych.

Branża interfejsów HMI – oprócz wszechobecnych trendów miniaturyzacji i redukcji poboru mocy – podlega także powszechnej ekspansji AI. O ile jednak w wielu przypadkach zaprzęgnięcie sztucznej inteligencji do realizacji określonych zadań bywa po prostu przerostem formy nad treścią albo zwyczajnym chwyt marketingowym, to w zakresie rozpoznawania trójwymiarowych gestów (na podstawie dość ograniczonych danych sensorycznych) użycie AI jest jak najbardziej na miejscu. Zwłaszcza, jeżeli uwzględnimy istotny fakt, o którym nie wolno zapominać projektantom interfejsów użytkownika: nawet najlepsza technologia stanie się źródłem irytacji, jeżeli nie zostanie dopracowana i solidnie przetestowana w najmniej sprzyjających warunkach otoczenia. Wszak czy ktokolwiek chciałby używać sprzętu, w którym obsługa za pomocą gestów przypominałaby trudną grę zręcznościową o niejasnych zasadach?

inż. Przemysław Musz, EP



Fotografia 18. Radar mikrofalowy BGT60TR13C marki Infineon (<https://t.ly/5Dyo2>)



Rysunek 13. Współpraca układu MXG3141 z kompatybilnymi kontrolerami dotykowych ekranów pojemnościowych (<https://t.ly/wlcyb>)

Pomiary charakterystyk częstotliwościowych (2)

Filtry aktywne m.cz. typu Sallen–Key



Pierwszy odcinek znajduje się pod adresem:
<https://ulubionykiosk.pl/media>

W kolejnej części cyklu poświęconego pomiarom charakterystyk częstotliwościowych liniowych układów przetwarzania sygnałów analogowych skupiono się na tzw. amplifiltrach m.cz., czyli filtrach aktywnych na małe częstotliwości. Element wzmacniający stanowi w nich nieodzowną część struktury układu. Niniejsze opracowanie stało się też okazją do przemycenia garści praktycznych rozwiązań projektowych. W publikacji pokazane i przeanalizowane zostały bowiem trzy projekty filtrów (dolnoprzepustowego, górnoprzepustowego oraz pasmowoprzepustowego), utworzone za pomocą środowiska projektowego firmy Analog Devices.

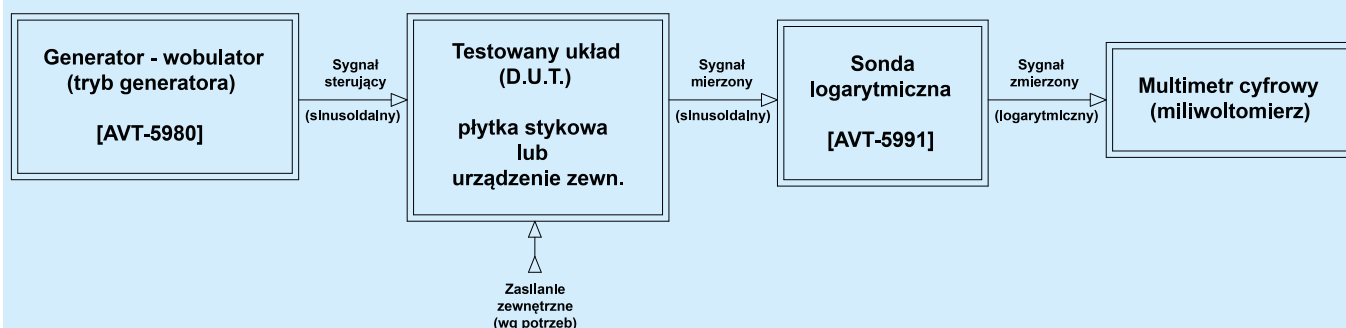
Wzmacniacze selektywne z założenia przeznaczone są do wydzielania i wzmacniania sygnałów w określonym pasmie częstotliwości oraz tłumienia sygnałów o częstotliwościach leżących poza wybranym fragmentem spektrum. W zależności od przyjętych wymagań mogą one mieć różny przebieg charakterystyki amplitudowej. Generalnie wyróżnia się wzmacniacze selektywne: dolnoprzepustowe, górnoprzepustowe, pasmowoprzepustowe oraz pasmowozaporowe, czyli selektywnie tłumiące sygnały. Wzmacniacze selektywne stosuje się m.in. w urządzeniach telekomunikacyjnych, audiowizualnych, radarowych, automatyce przemysłowej, technice pomiarowej, teledetekcji i w wielu innych dziedzinach elektroniki sygnałowej. Mogą one być wykonywane w różny sposób. Do dwóch podstawowych topologii należą: układy z obwodem rezonansowym, włączonym między dwa stopnie wzmacniające oraz topologie z obwodem selektywnym, włączonym w obwód sprzężenia zwrotnego. Wzmacniacze z selektywnym sprzężeniem zwrotnym nazywane są też filtrami aktywnymi lub amplifiltrami. Spośród licznych, stosowanych w praktyce koncepcji amplifiltrów, poczesne miejsce zajmują rozwiązania typu Sallen–Key, które wyróżnia prosta topologia, a także znaczna łatwość projektowania i implementacji w technologiach scalonych. Warto w tym miejscu zaznaczyć, że obwody selektywne w amplifiltrach m.cz. są w znakomitej większości oparte na elementach RC, które zapewniają znacznie bardziej zwartą konstrukcję i lepszą odporność na interakcje elektromagnetyczne z otoczeniem od komponentów indukcyjnych, które z kolei przy niższych częstotliwościach mogą też osiągać znaczne gabaryty i wysokie ceny. Z kolei we wzmacniaczach selektywnych w.cz., jako elementy kształtujące charakterystyki częstotliwościowe, zdecydowanie dominują elementy LC, piezoceramiczne i kwarcowe, które oferują znacznie korzystniejsze właściwości (m.in. dokładność, powtarzalność oraz stabilność temperaturową i czasową częstotliwości charakterystycznych, a także innych parametrów elektrycznych).



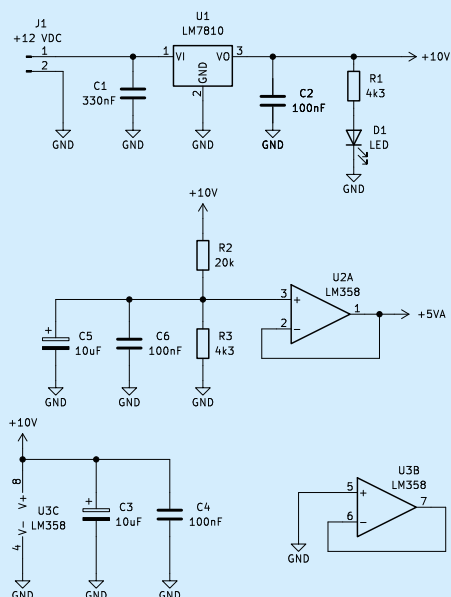
Metodyka pomiarów i zastosowane układy

W zaprezentowanych dalej testach jako źródło sygnału pomiarowego ponownie zastosowano generator-wobulator z modulem DDS [AVT5980, EP 4...6/2023], który pracował jako generator przestrajany ręcznie w wybranych zakresach częstotliwości. Schemat blokowy układu pomiarowego pokazano na **rysunku 1** – poziom sygnału na wyjściach amplifiltrów był odczytywany ręcznie za pośrednictwem szerokopasmowej sondy logarytmicznej AVT5991 (EP 8...9/2024) oraz multimetru cyfrowego (oscyloskop cyfrowy posłużył wyłącznie do kontroli kształtu mierzonego przebiegu napięcia). Dalej dane pomiarowe były wprowadzane do tabel, na podstawie których powstawały stosowne wykresy.

Na **rysunku 2** pokazano schematy układów zasilania oraz polaryzacji wejść użytych w testach wzmacniaczy operacyjnych, będących elementami aktywnymi w zaprojektowanych filtrach. **Rysunek 3** prezentuje okno startowe programu użytego do projektowania amplifiltrów. Jest to aplikacja sieciowa udostępniana bezpłatnie przez firmę Analog Devices (ADI) pod adresem internetowym



Rysunek 1. Schemat blokowy zastosowanego toru pomiarowego



Rysunek 2. Układy zasilania i polaryzacji wzmacniaczy operacyjnych

<https://tools.analog.com/en/filterwizard/>. We wszystkich trzech zaprojektowanych i zmierzonych układach filtrów aktywnych zastosowano popularny i łatwo dostępny, podwójny wzmacniacz operacyjny TL082, który w przedmiotowych projektach uznano za rozsądny zamiennik dla specjalizowanych, nowoczesnych wzmacniaczy operacyjnych produkcji ADI, rekomendowanych przez wspomniane narzędzie projektowe. **Tabela 1** przybliża wybrane parametry dynamiczne układu TL082. Wynika to, że układ oferuje bardzo wysoką impedancję wejściową i bardzo niski

Tabela 1. Wybrane parametry małosygnałowe wzmacniacza operacyjnego TL082

Symbol	Nazwa parametru	Typowa wartość	Jednostka
I _{ib}	Input bias current, T _{amb} =25°C	20	pA
A _{vd}	Large signal voltage gain, R _L =2 kΩ, V _o =±10 V, T _{amb} =25°C	200	V/mV
GBP	Gain bandwidth product, T _{amb} =25°C, V _{in} =10 mV, R _L =2 kΩ, C _L =100 pF, F=100 kHz	4	MHz
R _i	Input resistance	10 ¹²	Ω
THD	Total harmonic distortion, T _{amb} =25°C, F=1 kHz, R _L =2 kΩ, C _L =100 pF, A _v =20 dB, V _o =2 V _{pp}	0,01	%

wejściowy prąd polaryzacji – zapewne z uwagi na obecność tranzystorów JFET na wejściach różnicowych. W dodatku relatywnie wysokie wartości parametrów w opisywanych zastosowaniach: iloczyn „wzmocnienie-pasmo” GBP oraz wzmocnienia w otwartej pętli, pozwalają traktować ten podzespół niemalże jako idealny wzmacniacz operacyjny. Pozytywną percepcję zamiennego zastosowania tutaj układu TL082 dopełnia bardzo niski współczynnik zawartości harmonicznych THD w układzie testowym o dość znacznym wzmocnieniu i sporym napięciu wyjściowym.

Narzędzie projektowe firmy Analog Devices pozwalało na kilkuetapowe przygotowanie projektów filtrów – łącznie ze stosownymi symulacjami konkretnych rozwiązań układowych i dlatego w tym przypadku żadne odrębne symulacje nie były realizowane. Pełne cykle projektowe zostały opisane w dalszej części publikacji – wraz z opisami kolejnych etapów kreowania fizycznych układów.

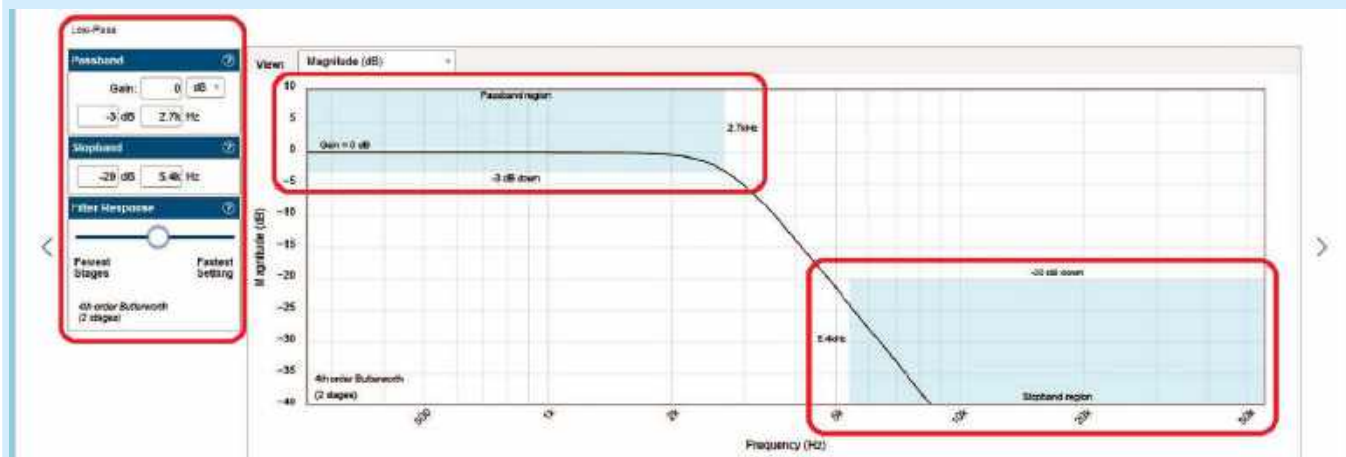
Dolnoprzepustowy amplifiltr m.cz. (LPF)

Jako punkt wyjścia do projektu wybrano filtr dolnoprzepustowy o częstotliwości granicznej $F_g=2,7$ kHz przy tłumieniu -3 dB. Mógłby to być np. filtr ograniczający pasmo odbieranego i/lub nadawanego sygnału głosowego w radioamatorskich emisjach SSB czy DSB. Po wybraniu w głównym oknie aplikacji projektowej (rysunek 3) stosownego prototypu filtru zostajemy przeniesieni do okna, w którym możemy wprowadzić podstawowe założenia do projektu (**rysunek 4**). Niezbyt wygórowane wymagania pod względem tłumienia w pasmie zaporowym (-20 dB przy $F=5,4$ kHz) pozwoliły uzyskać realizację filtru czwartego rzędu o charakterystyce maksymalnie płaskiej, tzn. wg tzw. aproksymacji Butterwortha. Suwak w dolnej części okna specyfikacji parametrów filtru (po lewej stronie ilustracji) pozwala na prosty wybór pomiędzy rozwiązaniem o bardziej stromej charakterystyce a takim, które wymaga bardziej złożonego rozwiązania projektowego (filtru wyższego rzędu).

Pokazane na **rysunku 5** kolejne okno aplikacji pokazuje już bazowy (wstępny) schemat projektowanego amplifiltru (widok „Circuit”). Na tym etapie można ustalić m.in.: dostępny zakres napięć zasilania wzmacniaczy operacyjnych oraz stopień swobody przy doborze wartości elementów RC w poszczególnych stopniach filtru, a także włączyć lub wyłączyć funkcję autokompensacji przez aplikację parametru GBW w modelu konkretnie



Rysunek 3. Okno startowe aplikacji do projektowania amplifiltrów (Analog Devices)



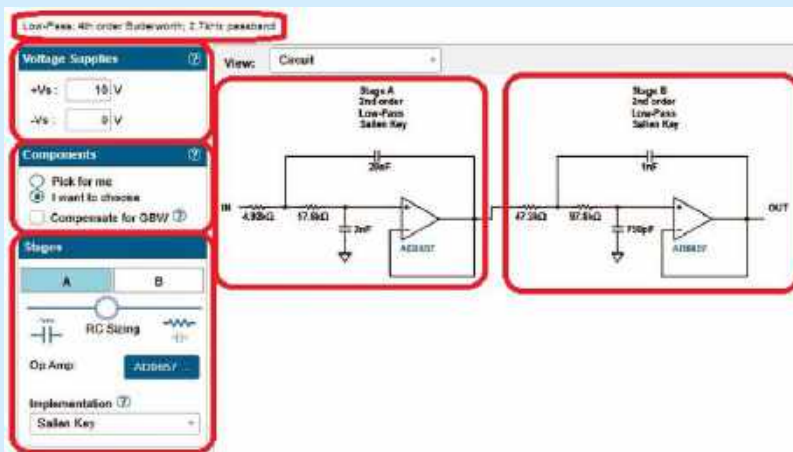
Rysunek 4. Zakożenia do projektu amplifiltru LPF typu Sallen-Key

wybranego (z obszernej listy sugerowanych podzespołów prod. firmy ADI) wzmacniacza operacyjnego. Autor projektu postanowił skorzystać z opcji ingerencji we wzajemne proporcje pomiędzy elementami RC w pierwszym stopniu filtru „A”, z uwagi na chęć uzyskania gwarantowanej impedancji wejściowej układu w całym interesującym pasmie częstotliwości na poziomie nie mniejszym niż około 100-krotność impedancji wyjściowej zastosowanego generatora, czyli około $100 \cdot 50 \Omega = 5 \text{ k}\Omega$. Jako topologię pojedynczego stopnia filtrującego drugiego rzędu wybrano architekturę typu Sallen-Key – nieco prostszą implementacyjnie od alternatywnie dostępnej topologii filtru z wielokrotnym sprzężeniem zwrotnym. W przypadku tej propozycji projektowej warto zwrócić uwagę na drobne niedociągnięcie, polegające na braku sugestii co do sposobu wprowadzenia do układu napięcia referencyjnego, polaryzującego wejście wzmacniaczy operacyjnych ($+V_s/2$). W praktyce zrobiono to na piechotę poprzez podanie wspomnianego napięcia referencyjnego wprost na wejście nieodwracające pierwszego wzmacniacza operacyjnego za pośrednictwem rezystora o wartości $1 \text{ M}\Omega$, którego znaczna wartość nie powinna zakłócić charakterystyki częstotliwościowej amplifiltru.

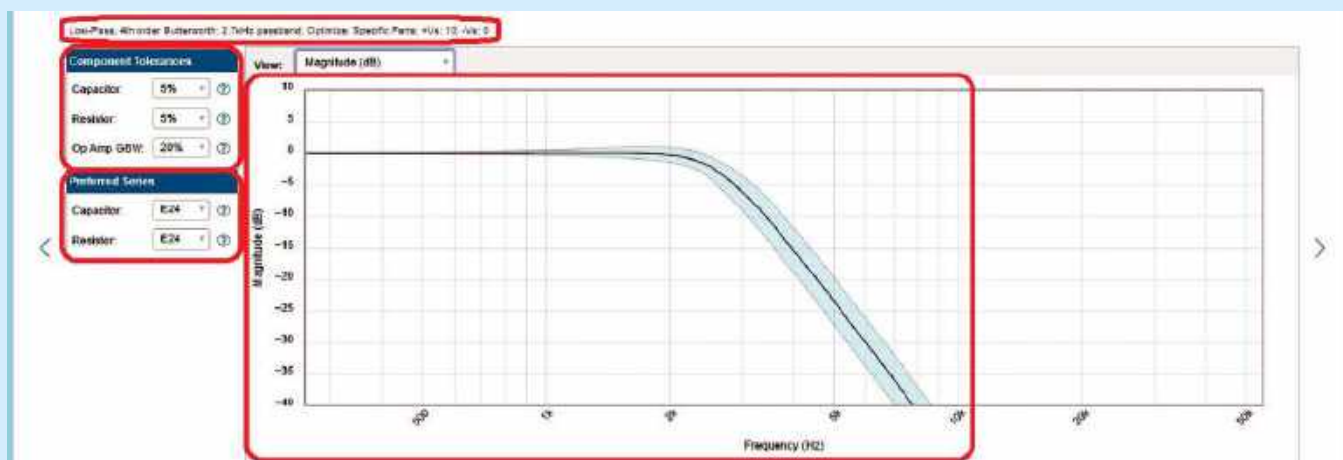
Rysunek 6 prezentuje kolejne okno projektu amplifiltru, w którym można wprowadzić dopuszczalne tolerancje kluczowych parametrów podzespołów w filtrze (wartości elementów RC oraz GBW wzmacniacza operacyjnego) a także typoszeręgi komponentów RC, z których aplikacja może w kolejnym kroku działania dobierać konkretne rezystancje i pojemności standardowe. Większą część

tego okna zajmuje wykres charakterystyki częstotliwościowej projektowanego układu (widok: „Magnitude (dB)”) wraz z wyróżnionymi błękitnym kolorem polami możliwych jej odstępstw od idealnego pierwowzoru, wynikającymi właśnie z tolerancji parametrów zastosowanych podzespołów.

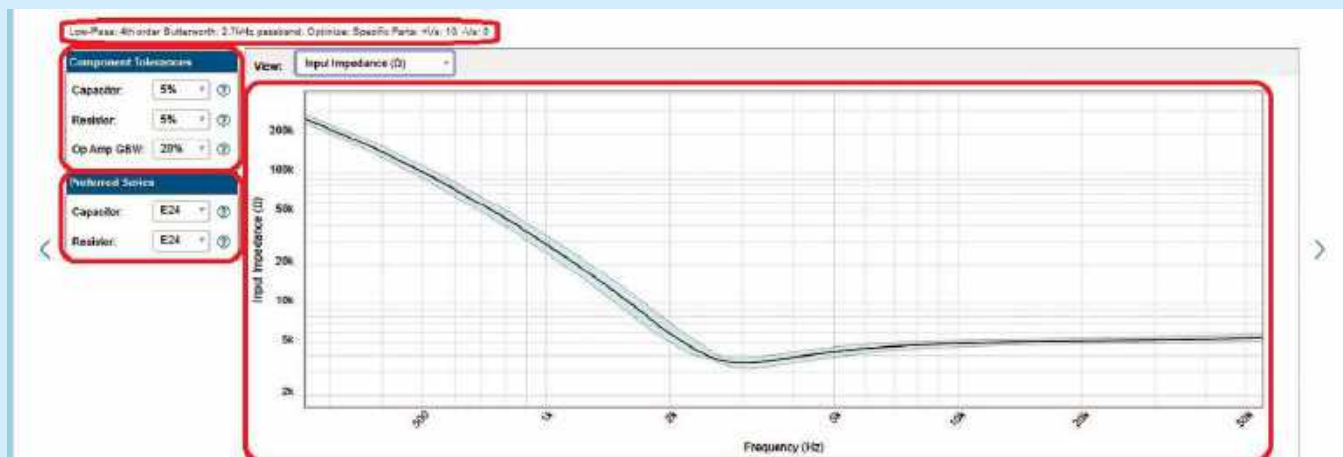
Rysunek 7 prezentuje widok „Input Impedance (Ω)” z wykresem impedancji wejściowej projektowanego amplifiltru. Jego analiza posłużyła do kontroli tego, czy zbyt niska impedancja wejściowa układu nie będzie w stanie zaburzyć pomiarów jego charakterystyki częstotliwościowej. Z kolei rysunek 8 (widok „Stages”) przybliży kluczowe własności poszczególnych stopni amplifiltru – ten tryb pracy może być przydatny szczególnie wtedy, gdy projektant nie jest do końca zadowolony



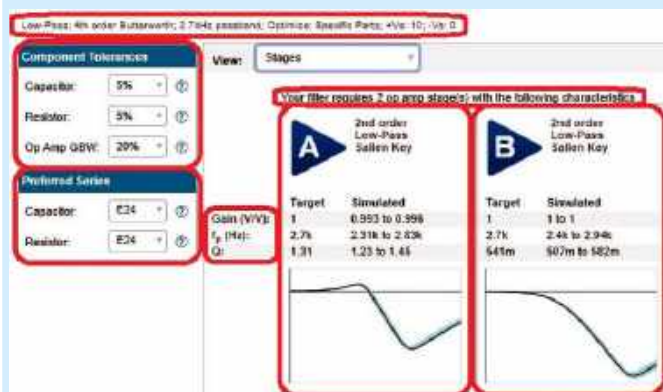
Rysunek 5. Bazowy schemat amplifiltru LPF typu Sallen-Key



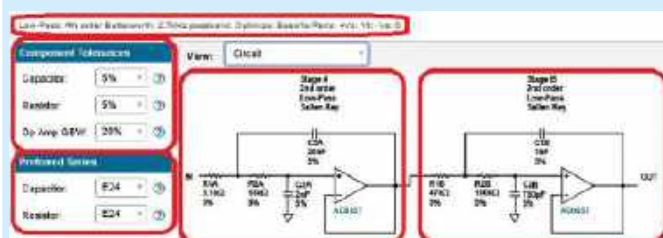
Rysunek 6. Charakterystyka częstotliwościowa amplifiltru LPF typu Sallen-Key



Rysunek 7. Impedancja wejściowa amplifiltru LPF typu Sallen-Key

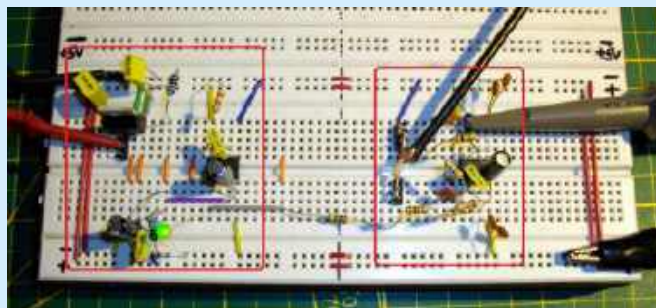


Rysunek 8. Własności poszczególnych stopni amplifiltru LPF typu Sallen–Key z uwzględnieniem tolerancji użytych elementów

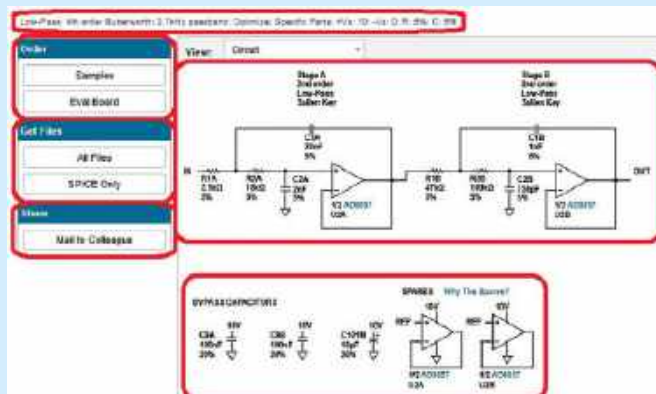


Rysunek 9. Ostateczny schemat amplifiltru LPF typu Sallen–Key

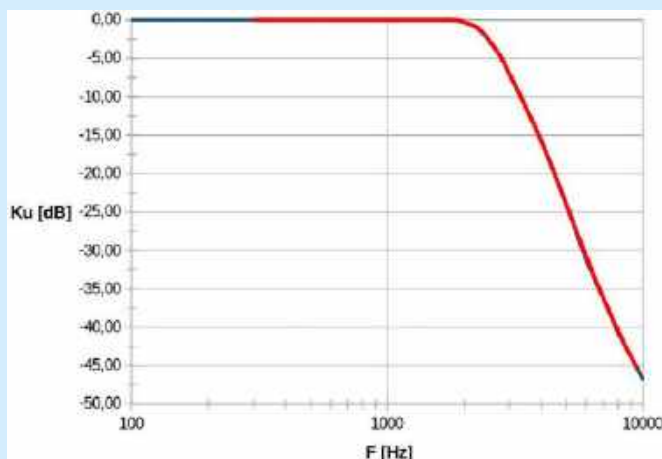
z uzyskanych własności całego filtra i poszukuje możliwości jego dalszej optymalizacji wyłącznie w ramach poszczególnych stopni. **Rysunek 9** (widok „Circuit”) prezentuje ostateczny schemat zaprojektowanego filtra aktywnego – już z zastosowaniem elementów RC o wskazanych dopuszczalnych tolerancjach i należących do preferowanych typoszeregów. **Rysunek 10** obejmuje pełne rozwiązanie projektowe – z uwzględnieniem pojemności blokujących zakłócenia zasilania poszczególnych wzmacniaczy operacyjnych czy sposobu zagospodarowania na firmowej płytce ewaluacyjnej pozostałych niewykorzystanych układów. W zasadzie taki schemat może już posłużyć do realizacji ostatecznej postaci amplifiltru. W lewej części tego okna można też znaleźć narzędzia wspierające zamówienie u producenta płytki ewaluacyjnej czy próbek wybranych wzmacniaczy operacyjnych, pobranie na dysk lokalny kompletu plików projektowych (w tym m.in. schematu, projektu PCB oraz BOM-u) lub tylko plików wsadowych do przeprowadzenia symulacji w środowiskach opartych na silniku SPICE. Dystrybutor aplikacji zapewnił też przycisk zachęcający do udostępnienia opracowanego projektu drogą mailową.



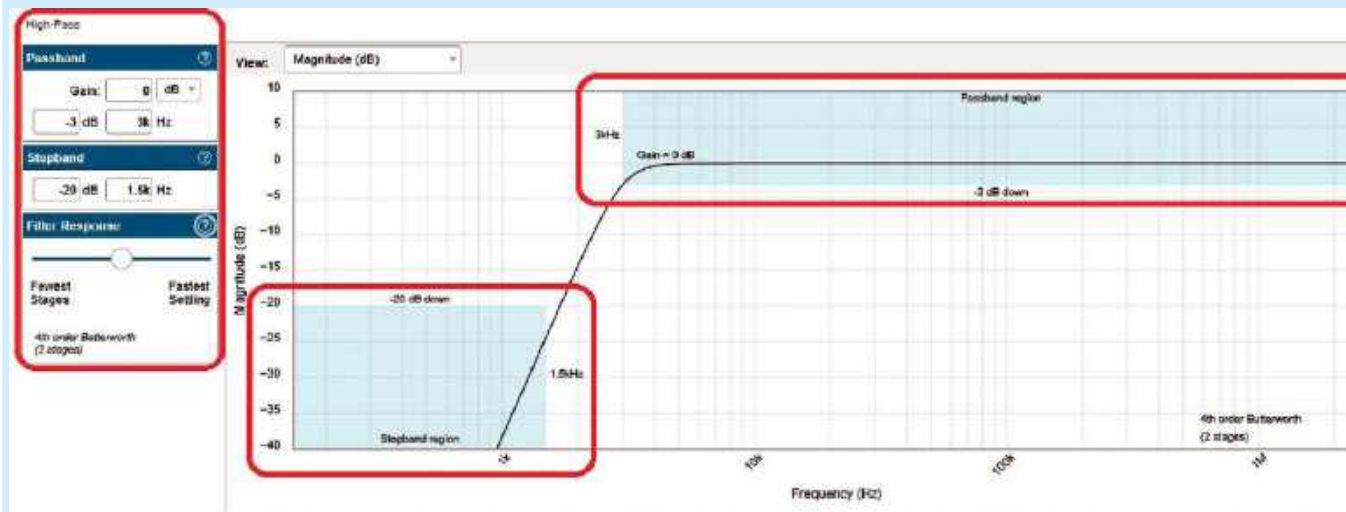
Fotografia 1. Roboczy model amplifiltru LPF typu Sallen–Key



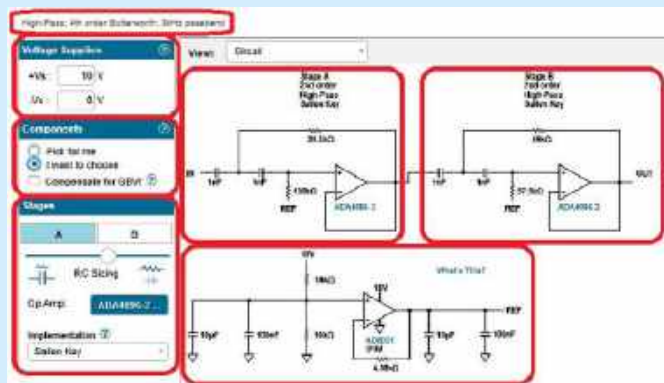
Rysunek 10. Sugestie odnośnie do realizacji ostatecznej postaci amplifiltru LPF typu Sallen-Key



Rysunek 11. Charakterystyka częstotliwościowa amplifiltru LPF typu Sallen-Key (pomiar)



Rysunek 12. Złożenia do projektu amplifiltru HPF typu Sallen-Key



Rysunek 13. Bazyowy schemat amplifiltru HPF typu Sallen-Key

Fotografia 1 pokazuje roboczy model zaprojektowanego uprzednio, dwustopniowego amplifiltru LPF czwartego rzędu. Tymczasowy model zbudowano na płytce stykowej, a jego konstrukcja (wbrew pozorom) wymagała znacznej staranności i ostrożności. W lewej części fotografii wyróżniono układy zasilania i polaryzacji, natomiast w części prawej widzimy właściwą implementację filtra aktywnego. Zagęszczenie częstotliwości pomiarowych dobrano stosownie do nieregularności charakterystyki amplitudowo-częstotliwościowej prototypu filtra. Liczbowe rezultaty pomiarów zostały ujęte w **tabeli 2** oraz zwizualizowane na wykresie na **rysunku 11**. Czerwona (wygładzona) linia trendu bardzo dobrze pokrywa się z niebieską linią odpowiadającą rzeczywistym

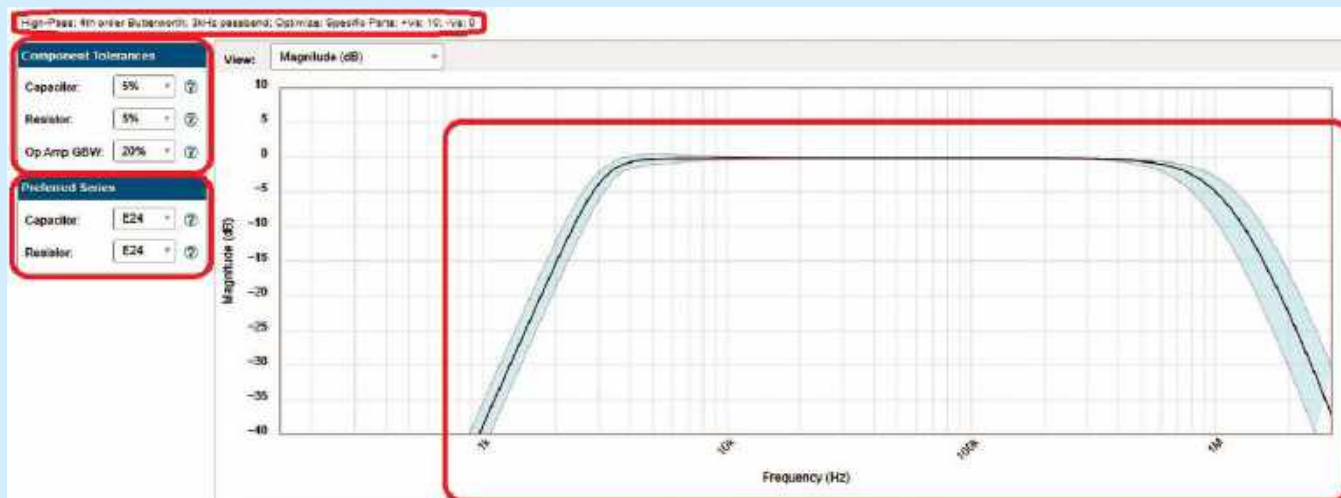
Tabela 2. Dane pomiarowe uzyskane dla amplifiltru LPF typu Sallen-Key

F [Hz]	Ku [dB]	F [Hz]	Ku [dB]	F [Hz]	Ku [dB]
100	0,00	2200	-0,83	4000	-15,72
500	0,00	2300	-1,27	4500	-20,00
1000	0,00	2400	-1,94	5000	-23,93
1500	0,00	2500	-2,66	6000	-31,29
1600	0,00	2600	-3,44	7000	-36,22
1700	0,00	2700	-4,18	8000	-40,83
1800	0,00	2800	-4,98	9000	-43,93
1900	-0,08	2900	-6,02	10000	-46,85
2000	-0,32	3000	-7,02		
2100	-0,57	3500	-11,58		

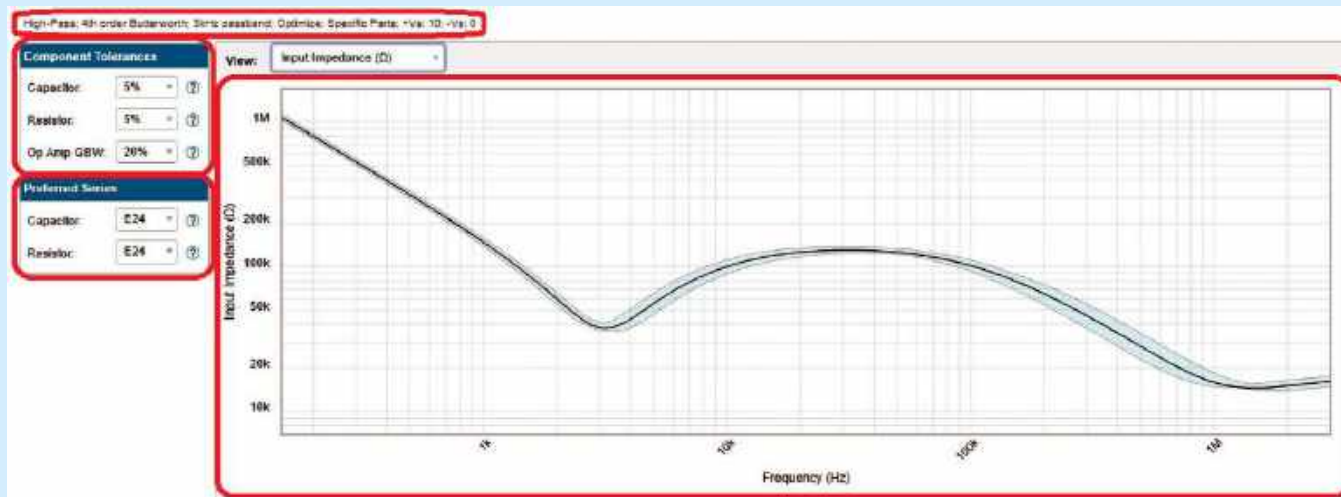
pomiarom. W dodatku obydwie krzywe bardzo dobrze zgadzają się z charakterystyką filtra prototypowego, co zasadniczo świadczy o poprawności realizacji zarówno projektu, jak i prototypu układu.

Górnoprzepustowy amplifiltr m.cz. (HPF)

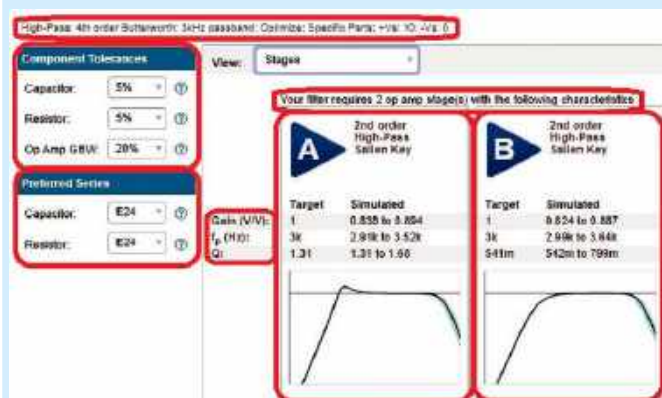
Przyczynkiem do podjęcia projektu tego filtra górnoprzepustowego pierwotnie była szersza koncepcja wzmacniacza stereofonicznego, wyłącznie na pasma górnych-średnich i wysokich częstotliwości akustycznych, który współpracowałby z dolno-średniotonowym, centralnym subwooferem. Przyjęto dolną częstotliwość graniczną filtra $F_d=3,0$ kHz przy tłumieniu -3 dB. Ponownie



Rysunek 14. Charakterystyka częstotliwościowa amplifiltru HPF typu Sallen-Key

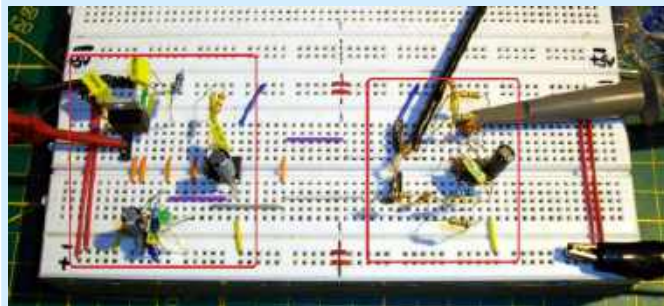


Rysunek 15. Impedancja wejściowa amplifiltru HPF typu Sallen-Key

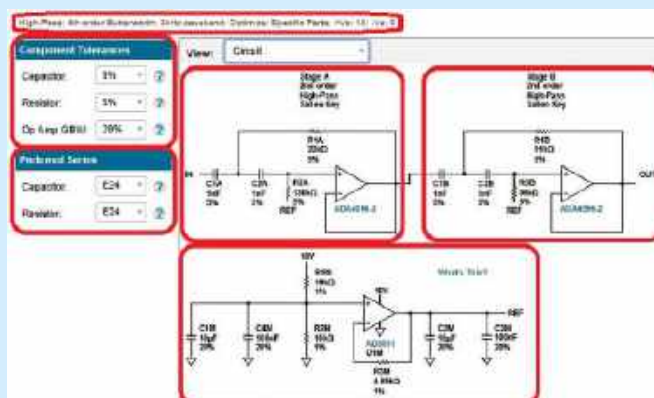


Rysunek 16. Właściwości poszczególnych stopni amplifiltru HPF typu Sallen-Key z uwzględnieniem tolerancji użytych elementów

najpierw w głównym oknie aplikacji projektowej (rysunek 3) wybieramy odpowiedni prototyp filtra, a następnie w kolejnym oknie wprowadzamy kluczowe założenia do projektu (rysunek 12). Umiarkowane wymagania odnośnie do tłumienia w pasmie zaporowym (-20 dB przy $F=1,5$ kHz) pozwoliły uzyskać realizację filtra czwartego rzędu o charakterystyce maksymalnie płaskiej (wg aproksymacji Butterwortha). Kolejne okno aplikacji (rysunek 13, widok „Circuit”) prezentuje wstępny schemat projektowanego filtra aktywnego. Ponownie skorzystano z możliwości ingerencji w proporcje między elementami RC w pierwszym stopniu filtra „A” w celu uzyskania odpowiedniej impedancji wejściowej układu w interesującym pasmie częstotliwości – nie mniejszej niż około 100 razy impedancja wyjściowa generatora, czyli 5 k Ω . W tym projekcie dla poszczególnych stopni także wybrano topologię typu Sallen-Key. Rysunek 14 pokazuje okno projektowe, przeznaczone do wprowadzania maksymalnych tolerancji kluczowych



Fotografia 2. Roboczy model amplifiltru HPF typu Sallen-Key



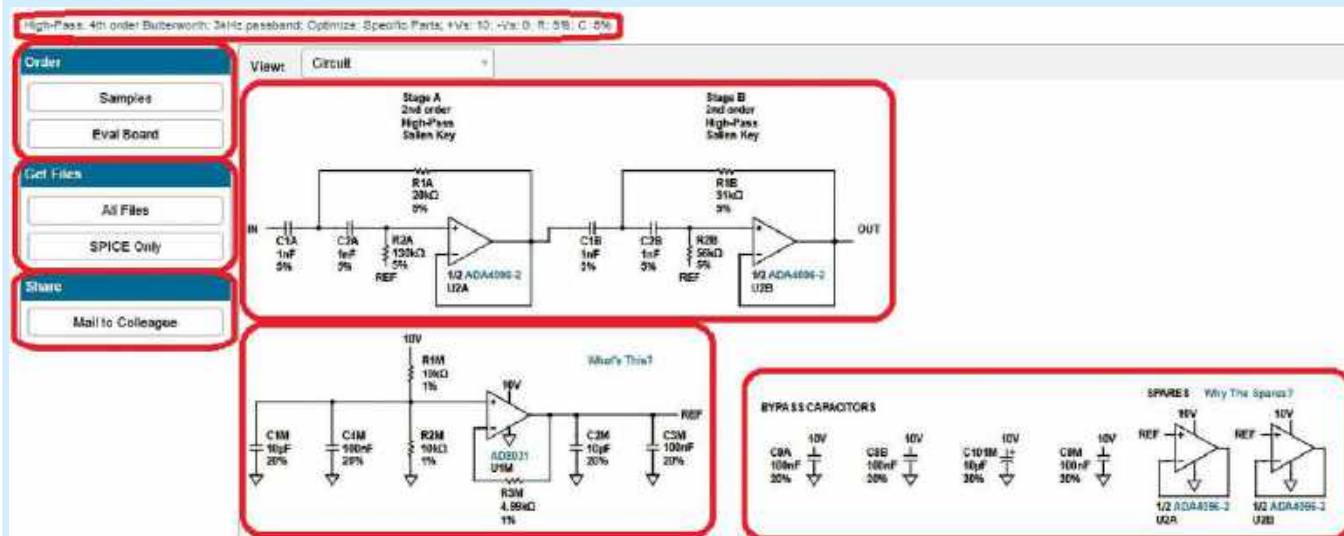
Rysunek 17. Ostateczny schemat amplifiltru HPF typu Sallen-Key

parametrów elementów filtra oraz preferowane typoszeregi elementów RC (do doboru konkretnych ich wartości) wraz z wynikiem wykresem opartym na analizie tolerancji.

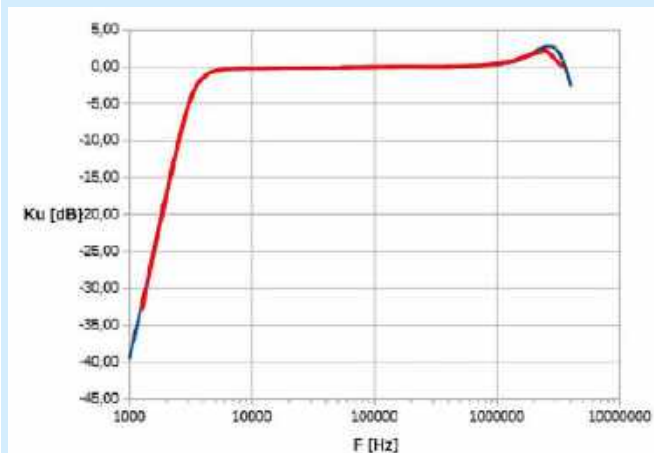
Rysunek 15 prezentuje widok „Input Impedance (Ω)” z wykresem impedancji wejściowej projektowanego filtra aktywnego. Wykazuje on spełnienie wstępnych założeń co do minimalnej wartości impedancji wejściowej układu. Natomiast rysunek 16 (widok „Stages”) pokazuje kluczowe parametry obu stopni amplifiltru na potrzeby ewentualnej dalszej optymalizacji poszczególnych stopni. Rysunek 17 (widok „Circuit”) pokazuje końcowy schemat układu (z zastosowaniem elementów RC

Tabela 3. Dane pomiarowe uzyskane dla amplifiltru HPF typu Sallen-Key

F [Hz]	Ku [dB]	F [Hz]	Ku [dB]	F [Hz]	Ku [dB]
1000	-39,57	3400	-3,08	6000	-0,30
1500	-26,51	3500	-2,67	8000	-0,22
2000	-16,97	3600	-2,28	10000	-0,22
2500	-10,21	3700	-2,00	100000	0,00
2600	-9,10	3800	-1,73	200000	0,00
2700	-8,12	3900	-1,60	500000	0,00
2800	-7,03	4000	-1,42	1000000	0,36
2900	-6,17	4200	-1,17	1500000	0,93
3000	-5,39	4400	-0,92	2000000	2,01
3100	-4,67	4600	-0,76	3000000	2,43
3200	-4,07	4800	-0,60	4000000	-2,67
3300	-3,50	5000	-0,53		



Rysunek 18. Sugestie odnośnie realizacji ostatecznej postaci amplifiltru HPF typu Sallen-Key



Rysunek 19. Charakterystyka częstotliwościowa amplifiltru HPF typu Sallen-Key (pomiar)

o dopuszczalnych wartościach), a **rysunek 18** ujmuje pełne rozwiązanie projektowe – np. do implementacji wprost na firmowej ewaluacyjnej PCB.

Fotografia 2 pokazuje roboczy model zaprojektowanego amplifiltru, zbudowany na płytce stykowej. Na fotografii wyodrębniono układy zasilania i polaryzacji oraz właściwą implementację filtru aktywnego. Zagęszczenie częstotliwości pomiarowych dostosowano do kształtu charakterystyki amplitudowo-częstotliwościowej prototypu filtru. Wyniki liczbowe pomiarów ujęto w **tabeli 3** oraz zaprezentowano graficznie na **rysunku 19**. Wyglądająca czerwona linia trendu bardzo dobrze pokrywa się z niebieską linią odpowiadającą rzeczywistym pomiarom, a obie krzywe również bardzo dobrze zgadzają się z charakterystyką filtru prototypowego – jednak

w okolicach wartości parametru GBW, zastosowanego w prototypie układu TL082 (ok. 4 MHz), można zaobserwować anomalię polegającą na miejscowym podbiciu charakterystyki przenoszenia filtru. Podbicie to wynika zapewne z braku w projekcie systemowej kompensacji parametrów wzmacniacza. Dodatkowo warto zwrócić uwagę na fakt, że amplifiltr, który z założenia miał być filtrem górnoprzepustowym, nie jest filtrem idealnym, ponieważ zarówno zaprojektowany komputerowo prototyp, jak i jego praktyczna realizacja mają ograniczone od góry pasmo przenoszenia. W tej sytuacji, w tym konkretnie zastosowaniu (filtr m.c.z. wydzielający pasmo wyższych średnich oraz wysokich częstotliwości akustycznych) powinien zostać profilaktycznie zaprojektowany jako filtr pasmowoprzepustowy (np. na pasmo 3...20 kHz), a być może także poprzedzony stosownym dolnoprzepustowym filtrem w pełni biernym (RC lub nawet LC – np. w topologii „ π ”).

Podsumowanie i wnioski

W materiale opisano pełne cykle projektowe oraz praktyczne pomiary charakterystyk częstotliwościowych dwóch aktywnych filtrów RC. Do przeprowadzenia projektów zastosowano profesjonalne środowisko programowe, oferowane bezpłatnie przez firmę Analog Devices w postaci aplikacji sieciowej. Zaprezentowane w artykule treści stały się m.in. pretekstem do przekazania garści wskazówek do projektowania i pomiarów podobnych filtrów aktywnych. W kolejnych odcinkach tego cyklu autor zamierza zaprezentować jeszcze dwie koncepcje analogowych amplifiltrów wąskopasmowych, projektowanych jednak całkowicie na piechotę, tzn. w oparciu o ogólnie znane topologie, stosowne wzory oraz powszechnie dostępne narzędzie symulacyjne, a następnie dokonać wprowadzenia do technologii DSP.

Adam Sobczyk, EP

REKLAMA

**świat
radio**
Magazyn wszystkich użytkowników eteru
KROTKOFALARSTWO CB · RADIOTECHNIKA

przejrzyj i kup na
www.ulubionykiosk.pl



Czytniki kodów kreskowych w praktyce (2)

Frontem do użytkownika: moduł skanera 1D/2D EM20-80 V2

W poprzedniej części kursu zapoznaliśmy się z budową, implementacją sprzętową oraz metodami konfiguracji miniaturowego silnika skanującego OEM typu EM3296V4, przeznaczonego przede wszystkim do urządzeń przenośnych i ubieralnych oraz niewielkich instalacji stacjonarnych. Tym razem omówimy diametralnie inny skaner kodów kreskowych 1D/2D, także opracowany przez firmę Newland – mowa o modelu EM20-80 V2.

Ergonomia ponad wszystko

Naprawdę trudno byłoby dziś znaleźć w Polsce osobę, która nie znałaby – chociażby z widzenia – wszechobecnych automatów paczkowych. Swoją niebywałą popularność zawdzięczają one łatwości odbierania przesyłek – niemal od samego początku urządzenia te były bowiem wyposażane we wbudowane skanery kodów QR, co znakomicie uprościło i przyspieszyło obsługę: wystarczy bowiem pokazać skanerowi kod (wyświetlony na ekranie smartfona lub ewentualnie wydrukowany na kartce papieru), zaś proces otwierania odpowiedniej skrytki zostanie automatycznie doprowadzony do końca.

W tego typu aplikacjach – a także we wszystkich innych zastosowaniach, w których użytkownik musi okazać maszynie odpowiedni kod kreskowy – szczególnie pożądane są rozwiązania redukujące efekt olśnienia (potocznie, choć niepoprawnie nazywanego także chwilowym oślepieniem) przez silne światło emitowane przez celownik i/lub oświetlacz skanera.

Konstrukcja silnika skanującego EM20-80 V2 (fotografie 1 i 2), będącego bohaterem tego odcinka naszego cyklu, już od samego początku procesu projektowego była przemyślana tak, by spełniać ten ważny – z punktu widzenia doświadczeń użytkownika – wymóg. Funkcję oświetlacza pełni tutaj bowiem zestaw diod LED, współpracujących z obszernym dyfuzorem, zainstalowanym dookoła wobec centralnie położonego obiektywu kamery CMOS. Światło



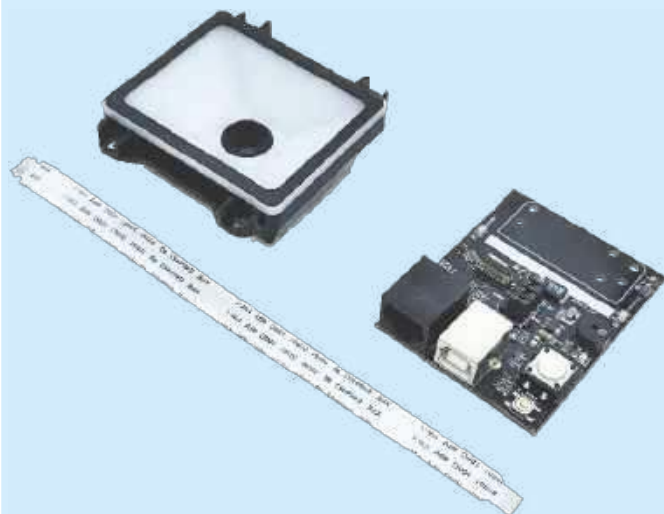
Pierwszy odcinek znajduje się pod adresem:
<https://ulubionykiosk.pl/media>

Więcej informacji:

Dystrybutorem skanerów OEM marki Newland na terenie Polski jest firma:
Kreski Spółka Jawna
ul. Mory 12, 01-303 Warszawa
<https://www.kreski.pl/>



emitowane przez tak zaprojektowany układ optyczny jest zarazem na tyle łagodne, by spojrzenie prosto w czytnik nie powodowało powstawania irytujących „blików” w oczach użytkownika oraz na tyle silne, by efektywnie iluminować okazany skanerowi kod kreskowy w niemal dowolnych warunkach oświetleniowych (fotografia 3). I to właśnie jest clou tytułowego podejścia frontem do użytkownika, co w nomenklaturze stosowanej przez markę Newland określa się mianem... *customer-facing*.



Fotografia 1. Zestaw złożony z silnika skanującego EM20-80 V2, płytki ewaluacyjnej oraz kompatybilnego przewodu FPC



Fotografia 2. Widok PCB skanera od tyłu



Fotografia 3. Moduł EM20-80 V2 w czasie skanowania wydrukowanego kodu QR

Mało tego – zastosowanie oświetlacza z dyfuzorem pozwala uniknąć odbić w sytuacjach, gdy kod jest odczytywany np. ze smartfona – zapisy przykładowych obrazów, zarejestrowanych przez kamerę czytnika EM20-80 V2 i zgranych do aplikacji EasySet, można zobaczyć na **rysunku 1**, zaś oryginalne obrazy monochromatyczne – bez żadnej dodatkowej obróbki – pokazano na **fotografiach 4...6**. I choć algorytmy skanerów korzystających z kolimowanej wiązki oświetlającej (takich jak np. moduł EM3296V4, opisany w poprzednim odcinku kursu) także doskonale radzą sobie ze skanowaniem kodów wyświetlanych na ekranach, to jednak w szczególnie niekorzystnych warunkach pracy brak przeszkadzających



Fotografia 4. Obraz zarejestrowany przez kamerę skanera podczas prezentacji kodu na matowym ekranie komputera



Fotografia 5. Obraz zarejestrowany przez kamerę skanera podczas prezentacji kodu wyświetlanego na ekranie smartfona



Rysunek 1. Widok okna programu EasySet w trybie przechwytywania nieprzetworzonego obrazu z kamery



Fotografia 6. Obraz zarejestrowany przez kamerę skanera podczas prezentacji kodu wydrukowanego na etykiecie termicznej

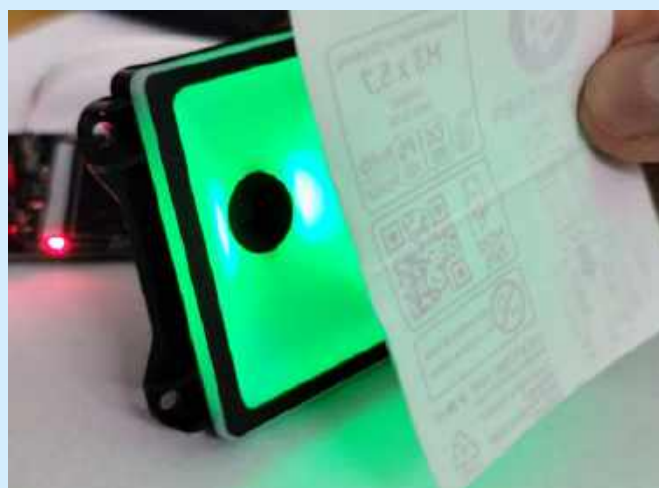
błysków zawsze będzie dodatkowym atutem, przemawiającym za implementacją czytnika typu *customer-facing*.

W tym miejscu należy dodać, że dyfuzor skanera EM20-80 V2 pełni także dodatkową funkcję – rozprasza bowiem światło zielonych diod LED, które sygnalizują dokonanie prawidłowego odczytu zaprezentowanego skanerowi kodu (**fotografia 7**). Jeżeli jednak szczególnie zależy nam na takim potwierdzeniu wizualnym, to zdecydowanie warto ustawić odpowiednio długi czas działania zielonych LED-ów – domyślnie wybrany czas (20 ms) jest bowiem na tyle krótki, że impuls w praktyce pozostaje niemal niezauważalny w typowych warunkach pracy. Nie jest to jednak żaden problem – oprogramowanie skanera umożliwia ustawienie dowolnej wartości z przedziału od 1 ms do 2500 ms – praktyczne testy wykazały, że optymalny czas załączenia diod zawiera się w przedziale od 300 ms do ok. 1 sekundy.

EM20-80 V2 w pigułce

Zanim przejdziemy do opisu implementacji modułu EM20-80 V2 w docelowym urządzeniu oraz konfiguracji czytnika za pomocą trzech różnych metod, przyjrzyjmy się najpierw najważniejszym parametrom technicznym naszego bohatera:

- dekodowanie: obsługa ok. 50 rodzajów kodów 1D i 2D,
- czujnikobrazu: kamera CMOS 640×480 px (monochromatyczna),
- oświetlacz: zestaw białych diod LED z dyfuzorem,
- sygnalizacja odczytu: zielone diody LED (za dyfuzorem) + wyjście do zewnętrznej diody LED + wyjście PWM do podłączenia buzzera,
- pole widzenia: 68° (poziomo) × 51° (pionowo),
- dopuszczalny kąt obrotu kodu w osi kamery: 360°,
- dopuszczalne nachylenie kodu w poziomie: ±45°,
- dopuszczalne nachylenie kodu w pionie: ±40°,



Fotografia 7. Zielone światło potwierdzające poprawny odczyt kodu

- minimalny kontrast wydruku: 30%,
- pobór prądu: 141 mA @ 3,3 V (w czasie skanowania), 93 mA @ 3,3 V (w trybie standby), 51 mA @ 5 V (w trybie standby),
- napięcie zasilania: 3,3...5,0 V (DC) $\pm 5\%$,
- oświetlenie zewnętrzne: 0...100 000 lx,
- zakres temperatur pracy: -40°C ... 60°C ,
- zakres temperatur przechowywania: -40°C ... 75°C ,
- wilgotność: 5%...95% (bez kondensacji),
- gwarancja: 2-letnia.

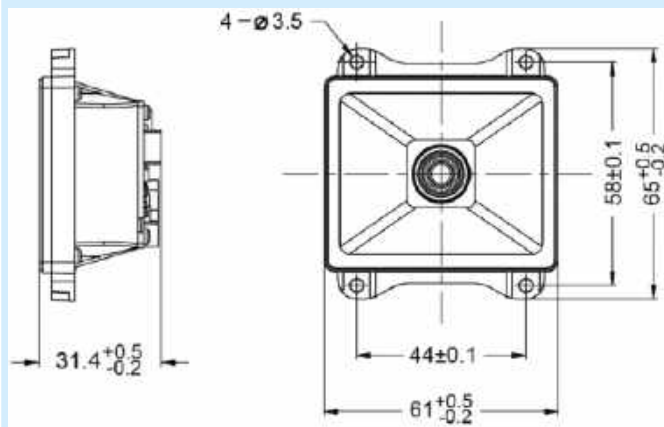
Wymiary, montaż mechaniczny

Moduł EM20-80 V2 jest zbudowany w oparciu o prostą (a dzięki temu niezawodną) konstrukcję mechaniczną – główna płyta drukowana, widoczna na fotografii 2, jest przykręcona do korpusu wykonanego z czarnego tworzywa sztucznego. Do płytki zamocowano także samą kamerę, której obiektyw został wyprowadzony w środkowej części dyfuzora, wytworzonego z mlecznobiałego tworzywa. Na jego zewnętrznym obwodzie umieszczono dodatkowo miękką uszczelkę – w trakcie integracji docelowego systemu należy tak zamontować silnik skanujący, by uszczelka pozostawała bezpośrednio docięnięta do okna optycznego – pozwoli to uniknąć jakichkolwiek niepożądanych odbić wewnętrznych oraz akumulacji zanieczyszczeń (kurzu, wilgoci czy małych owadów, mogących dostać się do wnętrza obudowy urządzenia). Do zamocowania modułu w docelowej obudowie przewidziano cztery otwory montażowe, rozmieszczone symetrycznie na wystających uszach i przystosowane do użycia ze śrubami M3 lub podobnymi. **Rysunek 2** prezentuje najważniejsze wymiary zewnętrzne czytnika.

Wybierając materiał osłonowy (okno optyczne), należy dopilnować, by jego powierzchnia była odpowiednio twarda i odporna na zarysowania. Producent zaleca zastosowanie szkła akrylowego (PMMA), tworzywa ADC (CR-39, czyli węglanu allilodiglikolu) bądź chemicznie wzmocnionego szkła. Minimalna transmitancja to 92%, zaś grubość powinna mieścić się w zakresie od 0,8 do 2,0 mm. Więcej szczegółów na temat parametrów materiałowych okna optycznego można znaleźć w dokumentacji [1].

Zasilanie, interfejsy elektryczne i pierwsze uruchomienie

Skaner EM20-80 V2 jest wyposażony w trzy złącza udostępniające zarówno linie zasilania, jak i interfejsy szeregowy oraz dodatkowe sygnały sterujące. Największe możliwości daje 12-pinowe gniazdo ZIF, współpracujące z taśmami FPC ze specjalnymi wycięciami blokującymi po obydwu stronach usztywniacza – podłączając przewód do tego gniazda, należy zatem zwrócić uwagę, czy wycięcia znalazły się za zaczepami, wykonanymi po obu stronach slotu w gnieździe (**fotografia 8**). Prawidłowo wykonane połączenia



Rysunek 2. Wymiary modułu EM20-80 V2



Fotografia 8. Prawidłowo umieszczona końcówka przewodu FPC w złączu ZIF modułu skanera. Montując przewód, należy zwrócić uwagę, czy wycięcia w usztywniaczu trafiły na specjalnie przewidziane do tego celu miejsce wewnątrz gniazda ZIF

między silnikiem skanującym a płytą ewaluacyjną pokazano na **fotografii 9**, zaś układ wyprowadzeń opisano w **tabeli 1**.

Pozostałe dwa złącza to niewielkie, 4-pinowe gniazda rastrowe typu 1.25T-4AWB marki Fuzhou Aoke Electronics Co. Ltd, o konstrukcji zbliżonej do gniazd z serii PicoBlade marki Molex – tutaj także mamy bowiem do czynienia z charakterystycznymi, płaskimi stykami, które na dodatek ustawione są w dokładnie takim samym rozstawie równym 1,25 mm. Pierwsze z nich – znajdujące się obok złącza ZIF – pozwala skorzystać z interfejsu RS-232 (**tabela 2**), zaś drugie (zamontowane na prostopadłym boku płytki drukowanej) udostępnia szynę USB (**tabela 3**). Dla maksymalnego ułatwienia implementacji w docelowym systemie producent wprowadził do sprzedaży gotowe kable USB (**fotografia 10**) oraz RS-232, pozwalające w prosty sposób podłączyć skaner do komputera bądź innego urządzenia nadrzędnego. W przypadku instalacji modułu w bezpośredniej bliskości płyty głównej lepszym wyborem będzie natomiast skorzystanie z gniazda ZIF i przewodu FPC – w takim

Tabela 1. Układ wyprowadzeń głównego złącza systemowego. I – wejście, O – wyjście push-pull, O_{od} – wyjście typu otwarty dren, PWR – zasilanie

Numer pinu złącza ZIF	Nazwa	Kierunek	Opis
1	nTRIG	I	wejście wyzwalające
2	nRESET	I	wejście zerujące
3	nGoodRead	O _{od}	wyjście sterujące wskaźnikiem LED (zewnętrznym)
4	nBEEPER	O _{od}	wyjście PWM do podłączenia buzzera
5	PWRDWN	–	niepodłączone
6	nRTS	O	żądanie nadawania (UART)
7	nCTS/USB D+	I (I/O)	gotowość do nadawania (UART)/linia danych USB
8	TXD	O	wyjście danych UART
9	RXD/USB D–	I (I/O)	wejście danych UART/linia danych USB
10	GND	PWR	masa
11	VIN	PWR	zasilanie 3,3...5,0 V
12	NC	–	niepodłączone



Fotografia 9. Kompletny zestaw ewaluacyjny gotowy do uruchomienia

Tabela 2. Układ wyprowadzeń złącza TTL-RS232

Numer pinu złącza UART (TTL RS-232)	Nazwa	Kierunek	Opis
1	VIN	PWR	zasilanie 3,3...5,0 V
2	RXD	I	wejście danych UART
3	TXD	O	wyjście danych UART
4	GND	PWR	masa

przypadku możemy bowiem zapewnić obsługę wszystkich dodatkowych linii, w tym wyprowadzeń sterujących zewnętrznym buzzerem, diodą LED, liniami kontrolującymi transmisję UART (RTS, CTS), a także obwodami resetu i zewnętrznego wyzwalania skanu.

Czytnik EM20-80 V2 jest domyślnie skonfigurowany w trybie emulacji urządzenia USB HID, a ściślej rzecz ujmując – klasycznej klawiatury komputerowej. Po podłączeniu kompletnego zestawu ewaluacyjnego do komputera PC za pomocą kabla USB z wtykiem typu B możemy od razu rozpocząć testy skanera – odczytanie kodu zostanie potwierdzone 20-milisekundowym błyskiem zielonych diod LED i 80-milisekundowym dźwiękiem

Tabela 3. Układ wyprowadzeń złącza USB

Numer pinu złącza USB	Nazwa	Kierunek	Opis
1	VIN	PWR	zasilanie 3,3...5,0 V
2	USB D-	I/O	linia danych USB
3	USB D+	I/O	linia danych USB
4	GND	PWR	masa

REKLAMA

Newland
SCANNING MADE SIMPLE

Odkryj więcej
newland-id.com



Dzięki naszym urządzeniom do skanowania codzienne zadania stają się prostsze, szybsze i bardziej intuicyjne.



Fotografia 10. Dedykowany przewód USB z 4-pinowym złączem rastrowym

buzzera o częstotliwości 2730 Hz. Poszczególne znaki, zapisane w zeskanowanym kodzie, zostaną przesłane jeden po drugim do komputera, bez dodatkowych opóźnień pomiędzy nimi. Istnieje oczywiście także możliwość podłączenia czytnika za pomocą wbudowanego portu USB oraz kabla pokazanego na fotografii 10 – wtedy jednak nie będziemy mogli skorzystać z udostępnianych przez zestaw ewaluacyjny, zewnętrznych elementów interfejsu użytkownika: przycisków resetowania i wyzwiania czy też buzzera i diody LED, potwierdzających dokonanie odczytu.

Konfiguracja za pomocą predefiniowanych kodów kreskowych

Podobnie jak w przypadku większości (jeśli nie wszystkich) współczesnych czytników kodów kreskowych, także skaner EM20-80 V2 daje możliwość konfiguracji setek parametrów i funkcji, udostępnianych przez oprogramowanie silnika skanującego, za pomocą predefiniowanych kodów „serwisowych”, których kompletną listę można znaleźć w dokumentacji [2].

Przykład 1. Zmiana częstotliwości buzzera

Oprogramowanie czytnika umożliwia ustawienie częstotliwości sygnału akustycznego (potwierdzającego dokonanie prawidłowego skanu) w zakresie od 20 Hz do 20 kHz. Założmy, że z uwagi na zastosowany rodzaj przetwornika chcemy zwiększyć wartość tego parametru z domyślnej 2,73 kHz na 4,2 kHz (czyli najwyższą predefiniowaną). W tym celu wskazujemy czytnikowi następujące kody, znajdujące się na stronie 14 dokumentacji [2]:

- **Enter Setup** – wejście w tryb programowania,
- **High (4200Hz)** – ustawienie właściwego parametru,
- **Exit Setup** – powrót do normalnego trybu pracy.

Aby sprawdzić, czy ustawienie zostało prawidłowo zapisane, wystarczy teraz zeskanować dowolny „zwykły” kod kreskowy 1D lub 2D (inny niż kody specjalne zawarte w [2]) – od razu da się zauważyć, że częstotliwość sygnału dźwiękowego uległa wyraźnemu podwyższeniu.

Przykład 2. Zmiana trybu skanowania

Domyślny tryb skanowania ustawiony w oprogramowaniu modułu to „Sense Mode” – czytnik monitoruje obraz rejestrowany przez kamerę i w razie relatywnie intensywnej zmiany (mogącej sugerować, że przed czytnikiem pojawił się kod) procesor załącza podświetlenie i dokonuje próby wykrycia kodu. W niektórych sytuacjach znacznie lepszym wyjściem jest skorzystanie z jednego z pozostałych trzech trybów pracy:

- **Level Mode** (tryb wyzwany poziomem logicznym) – ustawienie stanu aktywnego (niskiego) na wejściu wyzwajającym nTRIG powoduje rozpoczęcie sesji dekodowania, która trwa aż do momentu dokonania odczytu (wykrycia i zdekodowania kodu) lub do powrotu linii nTRIG w stan nieaktywny (wysoki).

- **Pulse Mode** (tryb impulsowy) – analogicznie jak w poprzednio opisanym trybie, jednak z tą różnicą, że rozpoczęcie sesji ma miejsce po chwilowej aktywacji, a następnie dezaktywacji linii nTRIG.
- **Continuous Mode** (tryb ciągły) – silnik skanujący automatycznie przemiata obraz z kamery w poszukiwaniu kodów i dekoduje je przez cały czas, niezależnie od tego, czy zawartość kodu zmieniła się, czy nie – powoduje to, że w czasie dłuższej prezentacji niezmiennego kodu (np. wydrukowanego lub statycznie wyświetlonego na ekranie smartfona) skaner będzie wielokrotnie wysyłał do urządzenia nadrzędnego ten sam ciąg znaków.

Dla przykładu zastosujemy tryb ciągły – w tym celu skanujemy kody znajdujące się na stronie 16 dokumentacji [2]:

- **Enter Setup**,
- **Continuous Mode**,
- **Exit Setup**.

Po 3 sekundach od zakończenia cyklu programowania czytnik włączy ciągłe podświetlenie i rozpocznie skanowanie.

Uwaga! Przed wykonaniem tego eksperymentu warto najpierw włączyć na komputerze dowolny edytor tekstowy (np. notatnik) i ustawić kursor na pustej stronie, gdyż ciągłe przysyłanie rozmaitych znaków w trybie emulacji klawiatury może spowodować niespodziewane zachowanie innych włączonych aplikacji (np. przeglądarki plików czy przeglądarki internetowej) – wynika to z prostego faktu: niektóre skróty klawiszowe mają postać... po jedynych liter, niewzbogaconych o klawisze specjalne (np. Shift, Ctrl czy Alt).

Włączenie trybu ciągłego skanowania przy pozostawionych domyślnych pozostałych ustawieniach może być dość niewygodne z uwagi na generowanie „ściany tekstu” w czasie prezentowania kodu czytnikowi. Aby szybko powrócić do bezpiecznych ustawień domyślnych, można posłużyć się kodem **Restore All Factory Defaults**, dostępnym na stronie 32 [2].

Obsługa programowania wsadowego

Nie ulega wątpliwości, że o ile zmiana pojedynczych ustawień za pomocą osobnych kodów serwisowych jest dość wygodna, o tyle już pełna konfiguracja skanera – odpowiednia do potrzeb docelowej aplikacji – byłaby niezwykle mozolna i pracochłonna. Dlatego też czytniki marki Newland są wyposażane w tzw. tryb programowania wsadowego (**Batch Programming**). Cała zabawa polega na tym, by przygotować odpowiedni ciąg znaków, łączący ze sobą wiele osobnych komend, a następnie wygenerować na jego podstawie dwuwymiarowy kod w formacie PDF417, QR Code lub Data Matrix – programowanie nastaw z jego użyciem będzie procesem nieporównanie szybszym i łatwiejszym niż osobne wpisywanie kolejnych poleceń i wertowanie prawie 300-stronicowej dokumentacji.

Składnia kodu wsadowego to w istocie ciąg znaków złożony z prefiksu („@”), po którym następują kolejne komendy rozdzielone średnikami (bez spacji!).

Założmy w tym momencie, że chcemy zmienić następujące ustawienia:

- czas świecenia zielonych diod LED po prawidłowym odczycie wydłużamy do 320 ms,
- dezaktywujemy buzzer (całkowicie wyłączamy sygnał potwierdzenia odczytu oraz melodyjkę odtwarzaną po włączeniu zasilania lub zresetowaniu modułu).

Jeżeli mamy już wybrane funkcje, które chcemy zmodyfikować, odszukujemy w dokumentacji odpowiednie kody serwisowe. Pod każdym z nich znajduje się alfanumeryczny ciąg znaków, rozpoczynający się od znaku „@” – i tak, dla zmiany czasu świecenia zielonych LED-ów jest to komenda:

- @GRLDUR320,

w przypadku wyłączenia sygnału dźwiękowego odtwarzanego po włączeniu modułu odnajdujemy polecenie:

- @PWBENA0

zaś wyłączenie sygnału potwierdzenia odczytu odbywa się przy użyciu ciągu znaków:

- @GRBENA0

Teraz wystarczy „skleić” ze sobą wszystkie trzy polecenia, pamiętając jedynie o tym, że znak „@” wpisujemy tylko raz na początku komendy. Otrzymujemy następujące polecenie wsadowe:

- @GRLDUR320;PWBENA0;GRBENA0

Teraz wystarczy już tylko utworzyć odpowiedni kod – w tym celu możemy skorzystać np. z darmowego generatora kodów QR, dostępnego online pod adresem: <https://qr.io> i widocznego na **rysunku 3**.

Uwaga – należy wybrać zakładkę „Text”, gdyż w innym przypadku do właściwego ciągu znaków zostaną dołączone dodatkowe znaki, które uniemożliwią poprawne działanie funkcji programowania wsadowego.

Po wpisaniu naszej komendy i odczekaniu krótkiej chwili, naszym oczom ukaże się gotowy kod QR, integrujący trzy wybrane przez nas polecenia.

Teraz wystarczy już tylko przejść na stronę 238 dokumentacji [2], a następnie zeskanować kody:

- **Enter Setup**,
- **Enable Batch Barcode**,
- **QR** (wygenerowany przed chwilą na stronie <https://qr.io>),
- **Exit Setup**.

Tryb wsadowy stanowi niebywale ułatwienie podczas programowania powtarzalnych konfiguracji, np. w produkcji seryjnej czy serwisie urządzeń. Co jednak począć, jeżeli urządzenie nadrzędne powinno móc samodzielnie zmieniać konfigurację modułu skanującego już na etapie eksploatacji? I tutaj właśnie dochodzimy do kolejnej funkcjonalności, oferowanej przez czytniki OEM marki Newland – programowania za pomocą interfejsu szeregowego UART (RS232) lub USB CDC.

Programowanie przez interfejs szeregowy

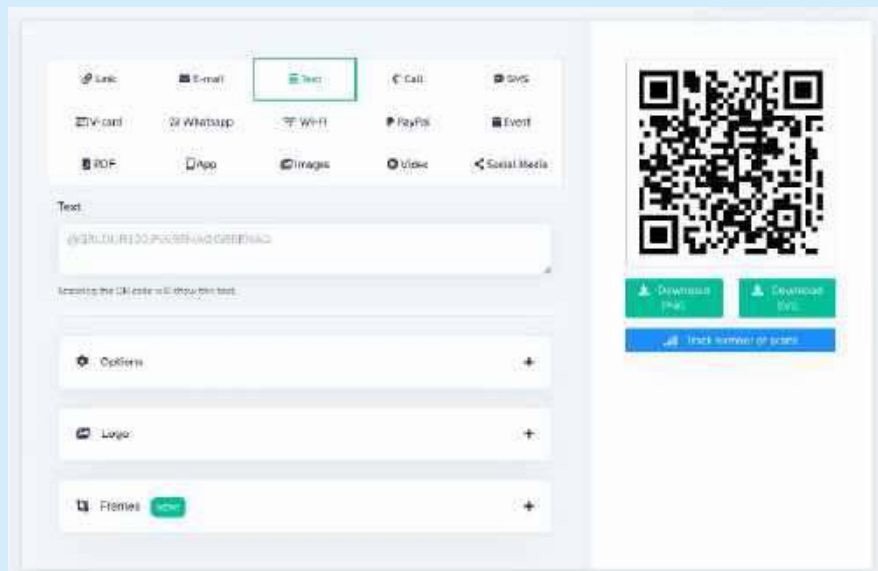
W poprzedniej części naszego kursu wspomnieliśmy już o możliwości przedstawienia trybu komunikacji czytników OEM, z czego korzystaliśmy zresztą w celu odczytu znaków za pomocą interfejsu szeregowego UART lub USB w trybie CDC (wirtualny port szeregowy). Teraz jednak nie będziemy odbierać, ale... wysyłać dane do naszego czytnika. W tym celu musimy skonfigurować go do odpowiedniego trybu pracy – najprościej będzie zastosować tryb **USB CDC**, który możemy włączyć kodem znajdującym się na stronie 63 dokumentacji [2].

Zanim jednak włączymy program terminalowy, by połączyć się z modułem i wysłać do niego komendę konfiguracyjną, zapoznajmy się pokrótce ze składnią poleceń, która w tym przypadku będzie nieco bardziej złożona niż w klasycznym trybie wsadowym z użyciem kodów programujących 2D. Powiedzmy, że chcemy przywrócić wyłączoną wcześniej obsługę buzzera.

Składnia podstawowej komendy programującej ma postać:

Prefix **Storage type** **Tag** **SubTag** **Data** **Suffix**

- **Prefix** – stały ciąg znaków o postaci „~<SOH>0000”, gdzie <SOH> to kod specjalny ASCII o nazwie *Start of Header* (0x01),



Rysunek 3. Kod QR do wsadowej konfiguracji, wygenerowany za pomocą aplikacji online dostępnej pod adresem: <https://qr.io>

- **Storage type** – określa tryb zapisu polecenia i przyjmuje dwie wartości: „@” (zapis do pamięci nieulotnej) lub „#” (zapis do pamięci RAM, tracony po resecie lub utracie napięcia zasilającego),
- **Tag** – trzyliterowy kod określający, do której grupy funkcjonalnej należy polecenie; w praktyce są to po prostu trzy pierwsze znaki „nazwy” danej komendy,
- **SubTag** – druga część nazwy polecenia, definiująca konkretną komendę z danej grupy,
- **Data** – dane do zapisu (np. wartość liczbowa),
- **Suffix** – stały ciąg znaków „;,<ETX>”.

Powyższa ramka, utworzona dla komendy @GRBENA1 (oznaczającej tymczasowe włączenie obsługi buzzera), przyjmuje następującą postać heksadecymalną:

7E 01 30 30 30 30 23 47 52 42 45 4E 41 31 3B 03

I to wszystko! Tak przygotowany ciąg znaków wystarczy przesłać za pomocą dowolnego programu terminalowego (obsługującego zapis szesnastkowy) do naszego modułu. Należy pamiętać, że domyślne nastawy interfejsu USB CDC to 9600 8N1, czyli 9600 baud, 8 bitów danych, brak obsługi parzystości i jeden bit stopu.

Prawidłowo zinterpretowana komenda zostanie natychmiast wdrożona, a potwierdzeniem tego faktu będzie zwrócenie przez czytnik nieco krótszego ciągu znaków – jego interpretację (w oparciu na dokumentacji [2]) pozostawiamy naszym Czytelnikom jako swego rodzaju zadanie domowe do samodzielnego wykonania.

Podsumowanie

W artykule zaprezentowaliśmy interesujący przykład nowoczesnego czytnika modułowego OEM o obszernym zakresie wbudowanych funkcjonalności i szerokich możliwościach konfiguracji. Warto bowiem wiedzieć, że oprócz programowania za pomocą kodów predefiniowanych i wsadowych oraz komend szeregowych, konfigurację modułu można zrealizować także za pośrednictwem oprogramowania EasySet, które opisaliśmy już w poprzednim odcinku niniejszego cyklu.

inż. Przemysław Musz, EP

[1] *EM20-80V2 OEM scan engine integration guide* [<https://t.ly/TwGdr>]

[2] *EM20-80V2 OEM scan engine user guide* [<https://t.ly/h3bWB>]



Implementacja systemu Linux na platformie STM32MP (2)

Pierwsza aplikacja



Pierwszy odcinek znajduje się pod adresem:
<https://ulubionykiosk.pl/media>

Graficzne interfejsy użytkownika są w dzisiejszych czasach nieodłącznym elementem wielu aplikacji. Często same urządzenia mają wbudowany wyświetlacz do komunikacji z użytkownikiem, zatem skorzystanie z systemu okienkowego – zintegrowanego z systemem Linux Embedded – wydaje się bardzo atrakcyjne.

W tej części kursu zajmiemy się napisaniem i uruchomieniem naszej pierwszej, prostej aplikacji wyposażonej w graficzny interfejs użytkownika, dalej nazywany skrótem GUI (ang. Graphical User Interface).

Nasz program zostanie napisany w języku Python, z uwagi na niski próg wejścia dla osób rozpoczynających swoją przygodę z Linuxem w systemach wbudowanych. W następnych odcinkach kursu będziemy się zajmowali pisaniem naszych aplikacji w językach takich jak C i C++. Wtedy również zaczniemy poznawać zasady tworzenia plików Makefile, niezbędnych do budowania aplikacji. Nie wybiegajmy jednak za nadto w przyszłość.

Hello World w Pythonie

Jak już wspomniałem na początku artykułu, w tym odcinku zajmiemy się napisaniem i uruchomieniem prostej aplikacji w Pythonie z obsługą GUI. Nie będziemy jednak musieli tutaj tworzyć żadnych plików (oczywiście poza głównym skryptem) w celu zbudowania naszej aplikacji.

Do opracowywania kodu używam środowiska VSCode z zainstalowaną wtyczką Pythona i – jak napisałem już w pierwszym odcinku naszego cyklu – pracuję na systemie Linux.

Kod zawarty na **listingu 1** jest krótki. Efektem jego działania będzie wyświetlenie na ekranie naszej płytki okienka aplikacji z przyciskiem „Hello World” (**rysunek 1**). Po naciśnięciu

przycisku w terminalu otrzymamy komunikat „Hello World!” (**rysunek 2**) a okno zostanie automatycznie zamknięte. Działanie aplikacji pokazano na **fotografiach 1 i 2**.

```
import gi

# Ustaw wymaganą wersję GTK na 3.0
gi.require_version('Gtk', '3.0')
from gi.repository import Gtk

# Definiujemy funkcję, która wypisuje "Hello World!" w konsoli, gdy zostanie wywołana
def print_hello(widget):
    print("Hello World!")

# Definiujemy główną funkcję, która inicjalizuje aplikację GUI
def activate(app):
    # Tworzymy główne okno aplikacji i przypisujemy je do aplikacji
    window = Gtk.ApplicationWindow(application=app)
    window.set_title("Window") # Ustawiamy tytuł okna
    window.set_default_size(200, 200) # Ustawiamy domyślny rozmiar okna

    # Tworzymy pudełko na przyciski i ustawiamy jego orientację na poziomą
    button_box = Gtk.ButtonBox.new(Gtk.Orientation.HORIZONTAL)
    window.add(button_box) # Dodajemy pudełko na przyciski do okna

    # Tworzymy przycisk z napisem "Hello World"
    button = Gtk.Button.new_with_label("Hello World")
    # łączymy zdarzenie "clicked" przycisku z funkcją print_hello
    button.connect("clicked", print_hello)
    # łączymy zdarzenie "clicked" przycisku z funkcją zamykającą okno
    button.connect("clicked", lambda w: window.destroy())
    # Dodajemy przycisk do pudełka na przyciski
    button_box.add(button)

    # Wyświetlamy okno i wszystkie jego elementy
    window.show_all()

# Główna funkcja, która inicjalizuje i uruchamia aplikację
def main():
    # Tworzymy nową aplikację GTK z unikalnym identyfikatorem aplikacji
    app = Gtk.Application.new("Demo.example", 0)
    # łączymy sygnał "activate" aplikacji z funkcją activate
    app.connect("activate", activate)
    # Uruchamiamy aplikację.
    app.run(None)

# Uruchamiamy główną funkcję, jeśli ten skrypt jest wykonywany bezpośrednio
if __name__ == "__main__":
    main()

Listing 1. Kod aplikacji
```

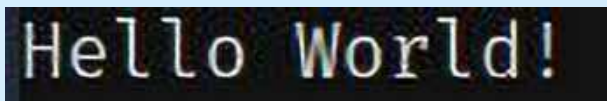
Wgrywanie kodu aplikacji na mikrokontroler

Aby wgranie naszej aplikacji na mikrokontroler było w ogóle możliwe, w pierwszej kolejności należy podpiąć zasilanie naszego zestawu ewaluacyjnego, a następnie podłączyć płytke do Internetu. Ważne jest, by płytka i komputer, na którym pracujemy, były podłączone z tą samą siecią LAN. W moim przypadku zastosowałem bezpośrednie połączenie płytki z routerem za pomocą patchcorda. Osobnym kablem mam podpięty komputer, również do tego samego routera.

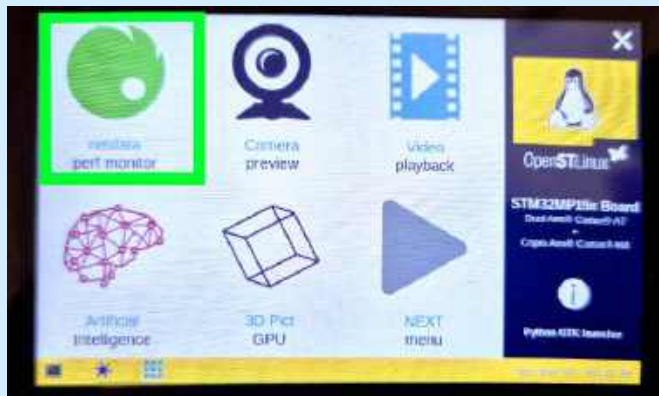
Aby uzyskać adres IP naszej płytki na ekranie głównym wyświetlanym przez płytke, wybieramy zieloną ikonkę „netdata” (**fotografia 3**). Po jej naciśnięciu pojawi się możliwość włączenia sieci Wi-Fi, jak również zostanie wyświetlony adres IP połączenia sieciowego, co widać na **fotografii 4**.



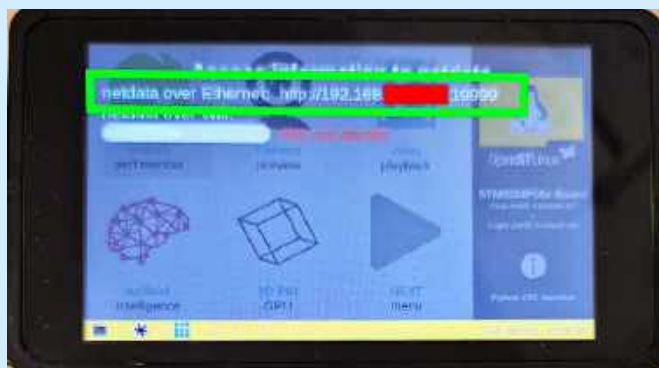
Rysunek 1. Efekt działania kodu z listingu 1



Rysunek 2. Komunikat wyświetlany w terminalu po naciśnięciu przycisku na formatce



Fotografia 1. Główny ekran systemu operacyjnego



Fotografia 2. Efekt działania aplikacji netdata

Za pomocą terminala zainstalowanego na głównym komputerze PC nawiązujemy połączenie z płytką za pomocą narzędzia ssh – stosujemy komendę (rysunek 1):

```
ssh root@<ip_płytki>
```

Po poprawnym nawiązaniu połączenia uzyskamy zdalny dostęp do terminalu linuksowego na płytce z poziomu komputera.

W tym momencie możemy wgrać kod naszej aplikacji na płytke. W tym celu otwieramy drugie okno terminalowe i – podając ścieżkę do miejsca, gdzie zapisaliśmy kod naszej aplikacji – za pomocą komendy:

```
scp 'nazwa_pliku' root@<ip_płytki>:/usr/local
```

przesyłamy plik do pamięci zestawu uruchomieniowego. Rysunek 2 pokazuje działanie komendy scp w terminalu. Użytkownicy systemów Linux korzystający z tmux nie muszą otwierać dodatkowego okna terminalu, wystarczy dodać nową sesję (**rysunek 3**).

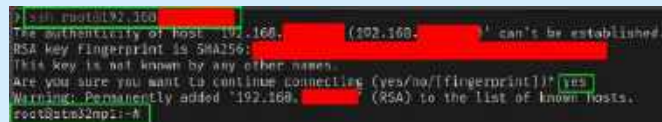
Uruchomienie naszej aplikacji

Korzystając z wcześniej nawiązanego połączenia ssh z naszą płytką, po poprawnym wgraniu kodu aplikacji, przyszedł nareszcie czas na jej uruchomienie. Używamy w tym celu komendy:

```
python3 /usr/local/'nazwa_pliku'.py
```

Jako nazwę pliku podajemy nazwę naszej aplikacji, w moim przypadku jest to po prostu main.py. Widok wpisanego polecenia można zobaczyć na **rysunku 4**. Po naciśnięciu klawisza Enter komenda zostanie wykonana – program uruchomi się, wyświetlając na ekranie płytki okienko, widoczne na fotografii 1.

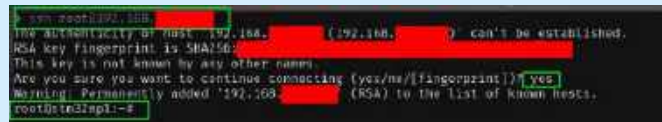
inż. Wiktor Hubaj



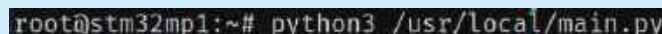
Rysunek 3. Połączenie z płytką przez ssh



Rysunek 4. Kopiowanie skryptu Pythona do pamięci płytki za pomocą polecenia scp



Rysunek 5. Okno programu tmux



Rysunek 6. Komenda uruchamiająca skrypt na płytce ewaluacyjnej

Kurs FPGA Lattice (27)

Obsługa monitora VGA



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

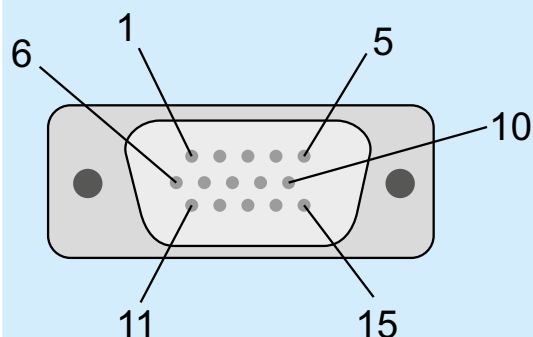
Przez wiele lat interfejs VGA dominował wśród sposobów łączenia monitorów z komputerami. Chociaż dziś już odchodzi do przeszłości, ustępując miejsca HDMI czy USB-C, to wciąż większość monitorów wspiera ten najważniejszy niegdyś standard. Warto wiedzieć, jak generować grafikę w FPGA za pomocą interfejsu VGA, bo ułatwi to opanowanie współczesniejszych rozwiązań w przyszłości.

Zacznijmy od omówienia sygnałów, jakie znajdują się w 15-pinowym konektorze D-Sub, używanym w interfejsie VGA. Zarówno monitor, jak i generator obrazu (tzn. karta graficzna komputera lub dowolne inne urządzenie dostarczające obrazy graficzne) wyposażone są w złącze żeńskie, a po obu stronach kabla powinny znajdować się wtyczki w wersji męskiej. Układ wyprowadzeń złącza pokazano na **rysunku 1**, a opis sygnałów zamieszczono w **tabeli 1**.

Linie te możemy podzielić na trzy grupy: sygnały analogowe odpowiedzialne za generowanie obrazu (piny 1, 2, 3, 13, 14), cyfrowa szyna danych służąca do identyfikacji monitora i jego możliwości oraz różne masy, które w praktyce można połączyć razem ze sobą.

Na przestrzeni lat zmieniały się standardy identyfikacji monitorów. W latach 90. wprowadzono standard EDID, bazujący na danych zapisanych w prostej pamięci EEPROM typu 24C z interfejsem I²C, która znajduje się wewnątrz monitora. Linie danych pamięci, tzn. SDA i SCL, wyprowadzone są wprost na złącze VGA. Pamięć ta działa nawet wtedy, gdy monitor nie jest podłączony do zasilania – w takiej sytuacji pamięć zasilana się z komputera za pośrednictwem piny 9, który dostarcza napięcie 5 V. Aby odczytać dane z pamięci, musimy wywołać ją na magistrali I²C spod adresu 0x50, a dalej postępować tak samo, jak w przypadku zwyczajnych pamięci szeregowych. W ten sposób można odczytać informacje o modelu monitora, producencie, obsługiwanych rozdzielczościach, kolorach, częstotliwości odświeżania obrazu, itp. Warto dodać, że w taki sam sposób rozpoznawane są monitory ze złączami DVI i HDMI. W tym odcinku kursu nie będziemy jednak szczegółowo omawiać tego tematu.

Piny 1, 2, 3 dostarczają informacje na temat składowych kolorystycznych: czerwonej, zielonej i niebieskiej aktualnie wyświetlanego piksela. Są to sterowane napięciowo wejścia analogowe. Zakres dopuszczalnych napięć może zmieniać się od 0 do 0,7 V – im wyższe napięcie, tym jaśniejszy kolor piksela. Każde z tych wejść połączone jest do masy poprzez rezystor 75 Ω. Nasz układ FPGA pracuje



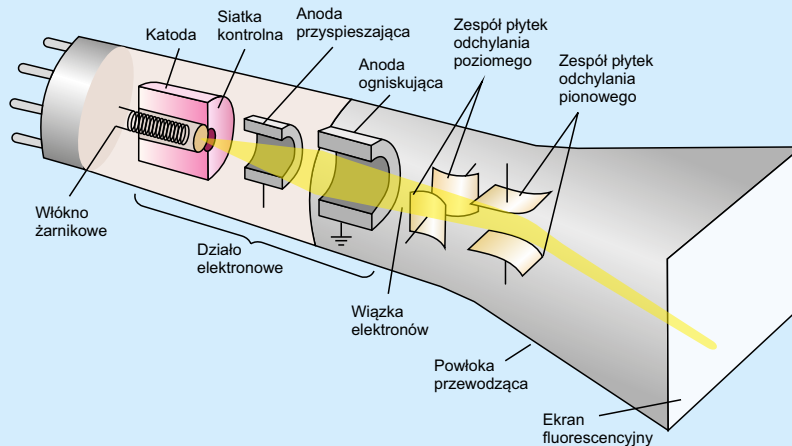
Rysunek 1. Pinout konektora VGA



Tabela 1. Opis sygnałów interfejsu VGA

Numer piny	Oznaczenie	Opis
1	RED	Analogowy sygnał koloru czerwonego
2	GREEN	Analogowy sygnał koloru zielonego
3	BLUE	Analogowy sygnał koloru niebieskiego
4	RES	Nie używany
5	GND	Masa
6	GND_RED	Masa koloru czerwonego
7	GND_GREEN	Masa koloru zielonego
8	GND_BLUE	Masa koloru niebieskiego
9	+5V	Zasilanie pamięci I ² C (EDID)
10	GND	Masa
11	RES	Nie używany
12	I2C_SDA	Identyfikacja monitora
13	HSYNC	Sygnał synchronizacji poziomej
14	VSNC	Sygnał synchronizacji pionowej
15	I2C_SCL	Identyfikacja monitora

przy napięciu 3,3 V, zatem pomiędzy nim a wyjściem VGA musimy umieścić dodatkowe rezystory, które wraz z rezystorami wewnątrz monitora będą tworzyć dzielnik napięcia. Z tego powodu na płycie User Interface Board zastosowano rezystory 270 Ω. Dzielnik napięcia zasilany z 3,3 V, składający się z rezystorów 270 Ω i 75 Ω, da bowiem na wyjściu napięcie 0,7 V. Proste i skuteczne.



Rysunek 2. Budowa lampy elektronopromieniowej w monitorze kineskopowym [3]

Tabela 2. Timing synchronizacji poziomej

Część cyklu	Piksele	Czas trwania [μs]	Napięcie na RGB [V]	Napięcie na HSYNC [V]
Obszar aktywny	640	25,422	0,0...0,7	3,3 lub 5,0
Front porch	16	0,636	0,0	3,3 lub 5,0
Sync pulse	96	3,813	0,0	0,0
Back porch	48	1,907	0,0	3,3 lub 5,0
Razem	800	31,778		

Tabela 3. Timing synchronizacji pionowej

Część cyklu	Linie	Czas trwania [ms]	Napięcie na RGB [V]	Napięcie na VSYNC [V]
Obszar aktywny	480	15,253	0,0...0,7	3,3 lub 5,0
Front porch	10	0,318	0,0	3,3 lub 5,0
Sync pulse	2	0,064	0,0	0,0
Back porch	33	1,049	0,0	3,3 lub 5,0
Razem	525	16,683		

Sygnały HSYNC i VSYNC służą do synchronizacji monitora z urządzeniem nadającym sygnał. Sposób synchronizacji jest bardzo staroświecki i wynika ze sposobu działania lamp elektronopromieniowych, zwanych w skrócie CRT. Budowę takiej lampy pokazano na **rysunku 2**. Składa się ona z działa elektronowego, które emituje wiązkę elektronów oraz z ekranu pokrytego luminoforem. Warstwa luminoforu świeci tym jaśniej, im więcej elektronów ją bombarduje. Ruchem elektronów sterują dwie pary elektrod kierujących, które – w zależności od przyłożonego napięcia – mogą wiązkę elektronów odchyłać pionowo i poziomo. W ten sposób wiązka przemiata cały ekran, linia po linii, zaczynając od lewego górnego rogu i kończąc na prawym dolnym.

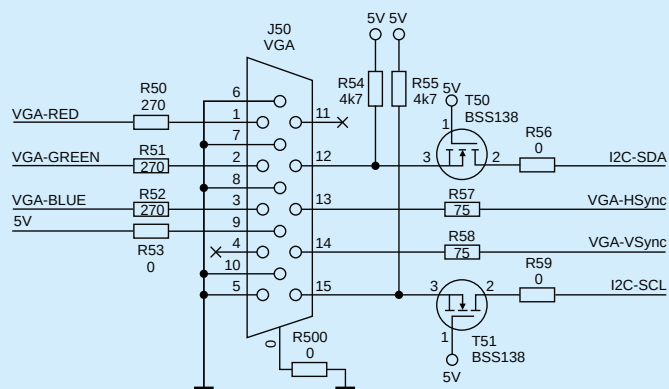
Sygnał synchronizacji pionowej mówi nam, że właśnie rozpoczyna się transmisja nowej klatki obrazu, a sygnał synchronizacji poziomej oznacza początek każdej kolejnej linii. Sygnały te nazywane są *sync pulse*. Jednak elektrody odchyłania poziomego i pionowego potrzebują trochę czasu, aby cofnąć wiązkę na pozycję początkową, czyli przeładować napięcie, jakie na nich występuje. Z tego powodu potrzebujemy dodatkowych dwóch faz cyklu, nazywanych *front porch* i *back porch*. W tym czasie obraz nie jest wyświetlany, a układ sterujący lampą CRT wytwarza odpowiednie napięcie, potrzebne do wyświetlenia kolejnej linii od początku.

W tabelach 2 i 3 zebrano czasy wszystkich czterech części cyklu pionowego i poziomego, obliczone dla rozdzielczości 640×480 pikseli i odświeżania 60 Hz. Jeżeli czas ten przeliczymy na piksele, wówczas okazuje się, że do monitora musimy dostarczyć obraz o rozdzielczości 800×525 pikseli. W tym nadmiarowym czasie sygnały odpowiedzialne za kolory czerwony, zielony i niebieski mają napięcie równe 0 V, a pracują jedynie linie synchronizacji pionowej i poziomej.

Policzmy pewną rzecz. Jeżeli zgodnie z tabelą 2 mamy wyświetlić 640 pikseli w czasie 25,422 μs, to czas wyświetlania jednego piksela powinien wynosić 39,722 ns. Jeżeli obliczymy odwrotność wyniku, to uzyskamy 25,175 MHz. Jest to częstotliwość, jaką najwygodniej będzie taktować układ FPGA, aby każdy takt zegarowy odpowiadał za wyświetlanie jednego piksela. Na szczęście monitory są tolerancyjne i dopuszczają dość duże odchyłki w częstotliwości taktowania sygnałów, zatem bez problemu można zastosować zegar o częstotliwości 25 MHz, bo jest popularniejszy i łatwiej dostępny.

Tabele z timingami dla innych rozdzielczości i częstotliwości odświeżania znajdziesz na stronie pod adresem [4].

Na **rysunku 3** zaprezentowano prosty schemat połączenia gniazda VGA z układem FPGA. Układ ten można zastosować



```
// Plik vga.v

`default_nettype none
module VGA(
    input wire Clock,      // Pin 25, musi być 25 MHz lub 25,175 MHz
    input wire Reset,      // Pin 17
    output reg Red_o,      // Pin 78
    output reg Green_o,    // Pin 10
    output reg Blue_o,     // Pin 9
    output reg HSync_o,    // Pin 1
    output reg VSync_o     // Pin 8
);

// Liczniki
reg [9:0] HCounter;      // 1
reg [9:0] VCounter;

// Maszyny stanów
reg [1:0] HState;        // 2
reg [1:0] VState;

// Stany
localparam ACTIVE = 0;   // 3
localparam FRONT = 1;
localparam SYNC = 2;
localparam BACK = 3;

// Logika licznika poziomego i pionowego
always @(posedge Clock, negedge Reset) begin // 4
    if(!Reset) begin
        HCounter <= 0;
        VCounter <= 0;
    end else begin
        if(HCounter != 799) // 5
            HCounter <= HCounter + 1'b1; // 6
        else begin
            HCounter <= 0; // 7
            if(VCounter != 524) // 8
                VCounter <= VCounter + 1'b1; // 9
            else // 10
                VCounter <= 0;
        end
    end
end

// Generowanie testowego kolorowego obrazu
wire [6:0] Sum = HCounter[9:4] + VCounter[9:4]; // 11
wire [2:0] ColorSelector = Sum[2:0]; // 12

// Pozioma maszyna stanów
always @(posedge Clock, negedge Reset) begin // 13
    if(!Reset) begin
        Red_o <= 0;
        Green_o <= 0;
        Blue_o <= 0;
        HSync_o <= 1; // 14
        HState <= ACTIVE;
    end

    else begin
        case(HState) // 15
            ACTIVE: begin // 16
                if(VState != ACTIVE) begin
                    {Red_o, Green_o, Blue_o} <= 3'b000;
                end else begin
                    case(ColorSelector) // 17
                        3'd0: {Red_o, Green_o, Blue_o} <= 3'b100; // czerwony
                        3'd1: {Red_o, Green_o, Blue_o} <= 3'b110; // żółty
                        3'd2: {Red_o, Green_o, Blue_o} <= 3'b010; // zielony
                        3'd3: {Red_o, Green_o, Blue_o} <= 3'b011; // błękitny
                        3'd4: {Red_o, Green_o, Blue_o} <= 3'b001; // niebieski
                        3'd5: {Red_o, Green_o, Blue_o} <= 3'b101; // fioletowy
                        3'd6: {Red_o, Green_o, Blue_o} <= 3'b000; // czarny
                        3'd7: {Red_o, Green_o, Blue_o} <= 3'b111; // biały
                    endcase
                end
            end
            FRONT: begin // 18
                HSync_o <= 1;
                if(HCounter == 639)
                    HState <= FRONT;
            end
            SYNC: begin // 19
                Red_o <= 0;
                Green_o <= 0;
                Blue_o <= 0;
                HSync_o <= 1;
                if(HCounter == 655)
                    HState <= SYNC;
            end
            BACK: begin // 20
                Red_o <= 0;
                Green_o <= 0;
                Blue_o <= 0;
                HSync_o <= 0;
                if(HCounter == 751)
                    HState <= BACK;
            end
        endcase
    end
end
```

wcześniejszych ćwiczeniach stosowaliśmy parametr **CLOCK_HZ** i moduły na ogół same sobie wyliczały różne dzielniki, ale tym razem będziemy pracować na pełnej częstotliwości sygnału zegarowego bez żadnych spowalniaczy. Projektując płytke MachXO2 Mega, umieściłem na niej generator kwarcowy o częstotliwości 25 MHz – właśnie na potrzeby dzisiejszego odcinka.

Kod modułu rozpoczynamy od zadeklarowania dwóch liczników **HCounter** i **VCounter** (linia 1). Ich zadaniem będzie wyznaczanie współrzędnych aktualnie wyświetlanego piksela. Licznik **HCounter** (H to skrót od *horizontal*) wyznacza, w której kolumnie znajduje się aktualnie wyświetlany piksel, a **VCounter** (*vertical*) określa numer wiersza. Licznik poziomy ma liczyć w zakresie od 0 do 799, a pionowy od 0 do 524. W obu przypadkach musimy użyć 10 bitów, aby pomieścić te liczby. Zwróć uwagę, że zliczać będziemy piksele leżące poza obszarem aktywnym, czyli liczniki wyznaczać będą także sygnały synchronizacji pionowej i poziomej.

W linii 2 i kolejnej tworzymy dwa rejestry maszyn stanów **HState** oraz **VState**, odpowiednio dla osi poziomej i pionowej. W obu przypadkach mamy tylko cztery fazy, zatem wystarczy, aby te rejestry były 2-bitowe. Nazwy stanów definiujemy w linii 3 i kolejnych.

W dalszej części kodu mamy trzy bloki `always`. Pierwszy z nich odpowiada za inkrementację liczników, drugi za oś poziomą, a trzeci – za pionową. Przejdźmy do omówienia pierwszego z nich (linia 4), którego logika jest bardzo prosta. Jeżeli licznik **HCounter** jest mniejszy od wartości maksymalnej, czyli 799 (linia 5), to inkrementujemy go (linia 6). Jeśli jednak obecna wartość licznika równa jest wartości maksymalnej, to zerujemy ów blok (linia 7) i wykonujemy analogiczną operację na liczniku **VCounter**. **VCounter** zwiększa swoją wartość (linia 9) dopiero wtedy, gdy licznik **HCounter** przepełnia się. Kiedy licznik **VCounter** osiągnie wartość maksymalną, czyli 524 (linia 8) i jednocześnie **HCounter** też osiąga maksimum, wówczas **VCounter** jest zerowany (linia 10).

Przeskoczmy teraz do kolejnego bloku `always`, który zaczyna się w linii 13. Jest on odpowiedzialny za generowanie sygnałów: czerwonego, zielonego, niebieskiego oraz synchronizacji poziomej **HSYNC**. Zwróć uwagę, że sygnał synchronizacji podczas spoczynku ma stan wysoki. Dlatego podczas resetu ustawiamy go na 1. Uważaj, by z przyzwyczajenia nie ustawić go na 0, tak jak w przypadku wszystkich innych sygnałów.

Omówimy najpierw sposób pracy maszyny stanów i generowanie sygnału **HSYNC**. Jest zrealizowany w sposób typowy – składa się z instrukcji `case` sprawdzającej wszystkie możliwe stany zmiennej **HState** (linia 15).

Po zresetowaniu układu rejestr stanów maszyny **HState** inicjalizowany jest wartością **ACTIVE**, czyli zerem. W takiej sytuacji

```

BACK: begin // 21
    Red_o <= 0;
    Green_o <= 0;
    Blue_o <= 0;
    HSync_o <= 1;
    if(HCounter == 799)
        HState <= ACTIVE;
    end
endcase
end
end

// Pionowa maszyna stanów
always @(posedge Clock, negedge Reset) begin // 22
    if(!Reset) begin
        VSync_o <= 1; // 23
        VState <= ACTIVE;
    end

    else if(HCounter == 799) begin // 24
        case(VState)
            ACTIVE: begin
                VSync_o <= 1;
                if(VCounter == 479)
                    VState <= FRONT;
            end

            FRONT: begin
                VSync_o <= 1;
                if(VCounter == 489)
                    VState <= SYNC;
            end

            SYNC: begin
                VSync_o <= 0;
                if(VCounter == 491)
                    VState <= BACK;
            end

            BACK: begin
                VSync_o <= 1;
                if(VCounter == 524)
                    VState <= ACTIVE;
            end
        endcase
    end
end

endmodule
`default_nettype wire

```

Listing 1. Kod pliku vga.v

Zobacz więcej:

- Repozytorium modułów wykorzystywanych w kursie <https://github.com/leonow32/verilog-fpga>
- Projekt w programie Diamond <https://tiny.pl/m6yb54bf>
- Budowa lampy elektronopromieniowej CRT https://tiny.pl/bryp_fc1
- Timing sygnałów VGA <http://tinyvga.com/vga-timing>

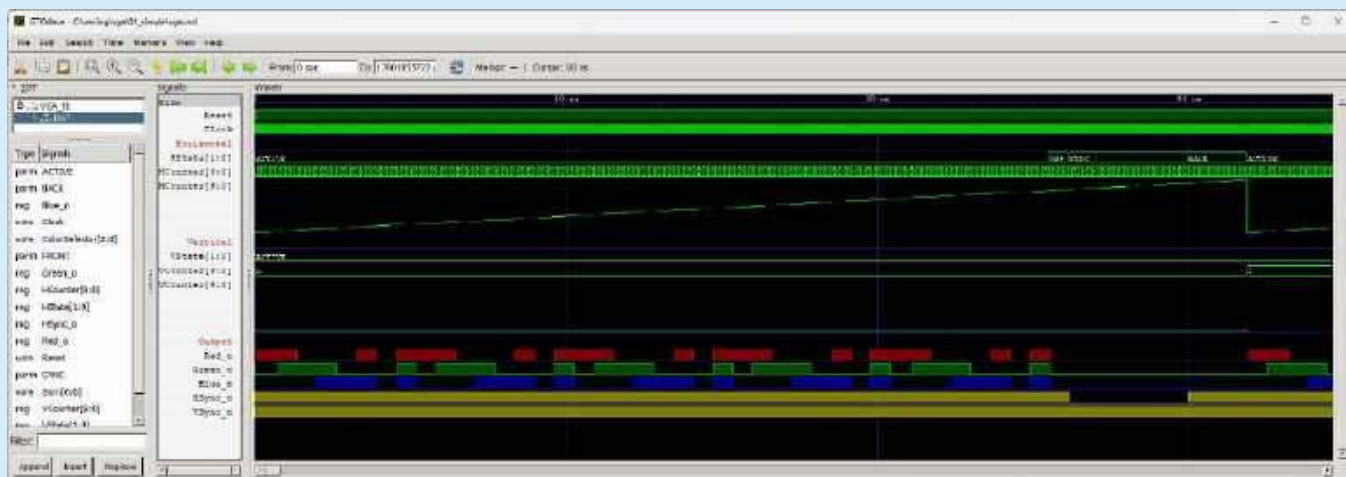
zajmujemy się generowaniem kolorów, ale tylko wtedy, kiedy pionowa maszyna stanów też pozostaje w fazie **ACTIVE**. Sprawdzamy to w linii 16. Jeżeli nie, to sygnały kolorów czerwonego, zielonego i niebieskiego powinny być w stanie niskim (sposób generowania barw omówimy kilka akapitów niżej). Niezależnie od tego wyjście **HSync_o** ma być w stanie wysokim. W taki sposób obsługujemy sygnały pierwszych 640 pikseli, tzn. od 0 do 639. W linii 18 sprawdzamy, czy właśnie teraz generowany jest ostatni piksel z tego zakresu i jeżeli tak, to wtedy do rejestru maszyny stanów wpisujemy wartość **FRONT**.

Następuje faza *front porch* (linia 19), w której sygnały kolorów są w stanie niskim, a sygnał synchronizacji cały czas pozostaje w stanie wysokim. Dzieje się to przez czas odpowiadający wyświetlaniu 16 pikseli. W tym czasie licznik **HCounter** cały czas zlicza takty zegarowe, zatem oczekujemy, aż doliczy do wartości 655 i wtedy do **HState** wpisujemy SYNC.

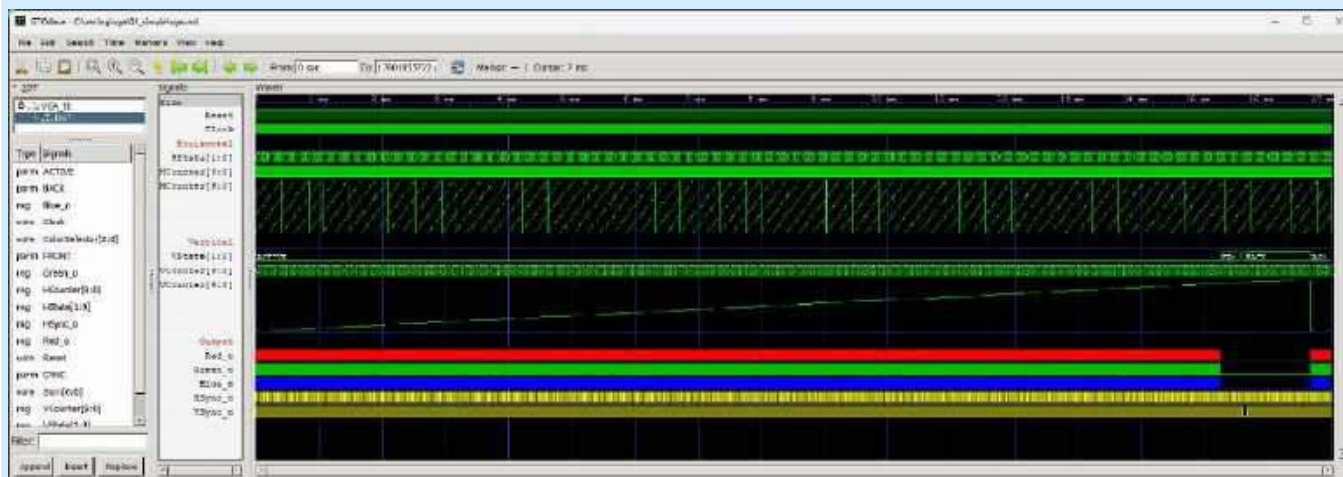
Faza synchronizacji jest bardzo podobna do poprzedniej, ale wyjście **HSync_o** przechodzi w stan niski, co stanowi właściwy sygnał synchronizacji poziomej. Trwa to tak długo, aż licznik **HCounter** osiągnie wartość 751 – wtedy przechodzimy do fazy *back porch* (linia 21). Tutaj również nie ma nic zaskakującego: wyjście **HSync_o** wraca w stan wysoki. Kiedy licznik **HCounter** osiągnie 799, zaczynamy wszystko od początku (licznik **HCounter** jest zerowany w poprzednim bloku *always*).

Zastanówmy się, w jaki sposób wygenerować sygnały odpowiedzialne za powstanie kolorowych kwadratów na monitorze. Chcemy, by kwadraty miały boki o długości 16 pikseli i były

w ośmiu podstawowych kolorach. Tak się dobrze składa, że rozdzielczość ekranu, tzn. 640×480 px, dzieli się przez 16. Zatem nasza mozaika składać się będzie z 40×30 kwadratów. Liczba 16 została wybrana nieprzypadkowo, ponieważ jest ona czwartą potęgą dwójki ($2^4=16$). Dzielenie przez liczbę będącą potęgą dwójki jest bardzo proste. Wystarczy „skreślić” tyle najmłodszych bitów dzielonej liczby, ile jest w wykładniku potęgi. Zatem aby podzielić liczbę przez 16, musimy usunąć cztery najmłodsze bity. Dzielić będziemy wartości liczników **HCounter** i **VCounter** (linia 11). Z tych zmiennych operatorem **[]** wybieramy wszystkie bity z wyjątkiem czterech



Rysunek 4. Zbliżenie na symulację pojedynczej linii



Rysunek 5. Symulacja całej jednej klatki obrazu

najmłodszych, po czym dodajemy je do siebie. Dostajemy wynik, który zmienia się w zakresie dużo większym, niż jest nam potrzebny. Dlatego w linii 12 tworzymy 3-bitową zmienną **ColorSelector**, do której przypisujemy tylko trzy najmłodsze bity sumy. W ten sposób otrzymujemy liczbę zwiększaną co 16 pikseli i zmieniającą się w zakresie od 0 do 7. Ponadto w pierwszej linii kwadratów zaczynamy od 0, a w drugiej linii od 1, w trzeciej od 2 i tak dalej.

Jak zapewne się domyślasz, zawartość zmiennej **ColorSelector** ma decydować o tym, który kolor jest obecnie wyświetlany. Zobacz linię 17, zawierającą kolejną instrukcję case. Wykonuje się ona tylko

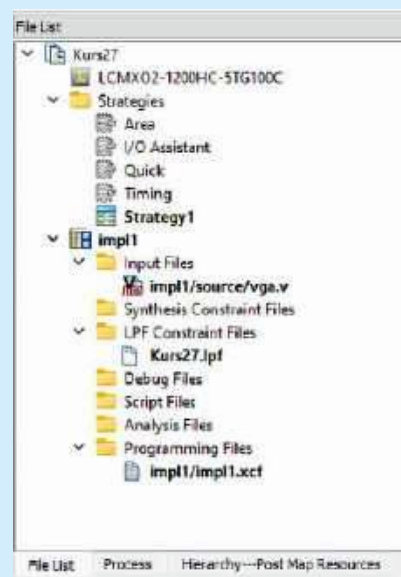
wtedy, gdy obie maszyny stanów są w fazie aktywnej. Na podstawie wartości zmiennej **ColorSelector** następuje przypisanie wyjść **Red_o**, **Green_o**, **Blue_o**. Aby kod był bardziej zwięzły, przypisujemy wartość wszystkim trzem zmiennym jednocześnie za pomocą operatora konkatencji {}.

Do omówienia pozostaje już tylko trzeci blok always (linia 22), odpowiedzialny za pionową maszynę stanów i sygnał synchronizacji pionowej VSYNC. Sygnał ten jest w stanie wysokim i dlatego podczas resetu wpisujemy do niego 1 (linia 23). Dalej mamy maszynę stanów, jednak jest ona przełączana dopiero wtedy, gdy licznik **HCounter** osiąga wartość maksymalną (linia 24). Konstrukcja pionowej maszyny stanów jest bardzo podobna do maszyny poziomej, więc nie będziemy jej dokładnie omawiać.

Testbench modułu VGA

Testbench naszego modułu VGA będzie nieskomplikowany. Jego kod pokazano na **listingu 2** – ogranicza się do utworzenia instancji testowanego modułu, puszczeniu go w ruch i czekaniu, aż wygeneruje całą klatkę obrazu. Koniec klatki rozpoznajemy w linii 1, gdzie sprawdzamy, czy oba liczniki jednocześnie mają wartość maksymalną. Dalej, w linii 2 czekamy jeszcze trochę czasu, aby ponownie zobaczyć początek klatki.

Aby uruchomić symulację w Icarus Verilog, należy napisać prosty skrypt, którego kod zaprezentowano na **listingu 3**.



Rysunek 6. Drzewo projektu w Lattice Diamond

```
// Plik vga_tb.v

`timescale 1ns/1ps
`default_nettype none
module VGA_tb();

    parameter CLOCK_HZ = 25_175_000;

    // Generator sygnału zegarowego
    reg Clock = 1'b1;
    always begin
        #(1_000_000_000.0 / (2.0 * CLOCK_HZ));
        Clock = !Clock;
    end

    // Zmienne
    reg Reset = 1'b0;

    // Instancja testowanego modułu
    VGA DUT(
        .Clock(Clock),
        .Reset(Reset),
        .HSync_o(),
        .VSync_o(),
        .Red_o(),
        .Green_o(),
        .Blue_o()
    );

    // Eksport wyników symulacji
    initial begin
        $dumpfile("vga.vcd");
        $dumpvars(0, VGA_tb);
    end

    // Sekwencja testowa
    initial begin
        $timeformat(-6, 3, "us", 12);
        $display("===== START =====");

        @(posedge Clock);
        Reset <= 1'b1;

        wait(DUT.VCounter == 524 && DUT.HCounter == 799); // 1
        wait(DUT.VCounter == 10); // 2

        $display("===== END =====");
        $finish;
    end

endmodule
`default_nettype wire

Listing 2. Kod pliku vga_tb.v
```

```
@echo off
iverilog -o vga.o ^
    vga.v ^
    vga_tb.v
vvp vga.o
del vga.o
```

Listing 3. Kod pliku vga.bat

Otwórz plik vga.vcd w przeglądarce GTKWave i skonfiguruj ją tak, by uzyskać efekt pokazany na **rysunkach 4 i 5**. Pierwszy z nich reprezentuje zbliżenie na symulację jednej linii obrazu. Widzimy, jak pracuje pozioma maszyna stanów, korzystająca ze zmiennej **HState** oraz w jaki sposób generowane są sygnały na wyjściach **Red_o**, **Green_o** i **Blue_o**. Ponadto widzimy, w jaki sposób zmienia się sygnał synchronizacji poziomej. Rysunek 5 pokazuje natomiast symulację całej klatki obrazu o rozdzielczości 640×480 pikseli – tutaj możemy analizować pracę pionowej maszyny stanów ze zmienną **VState** oraz sposób, w jaki zmienia się sygnał synchronizacji pionowej. Za pomocą kursorów można zmierzyć czas trwania wszystkich sygnałów, aby zweryfikować, czy są one zgodne z tym, co podano wcześniej w tabelach 2 i 3.

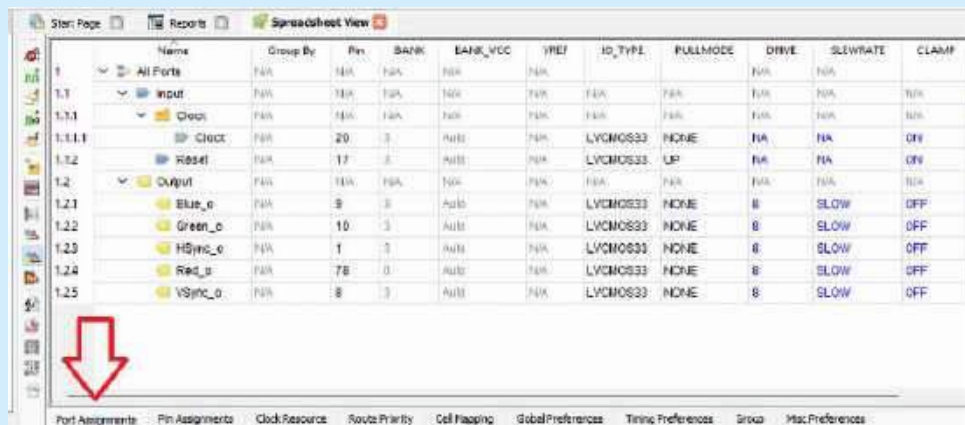
Testy w FPGA

W tym odcinku nie będzie modułu **top**. Nasz projekt jest na tyle prosty, że moduł VGA może bezpośrednio pełnić funkcję modułu nadrzędnego. Z tego powodu drzewko projektu w Lattice Diamond powinno wyglądać tak, jak pokazano na **rysunku 6**.

Przeprowadź syntezę, a następnie otwórz narzędzie Spreadsheet i skonfiguruj piny układu FPGA w sposób pokazany na **rysunku 7**. Pamiętaj, że używamy zewnętrznego generatora sygnału zegarowego, więc powinniśmy także podać jego częstotliwość w Timing Preferences (zakładki na dole Spreadsheet), co pokazuje **rysunek 8**.

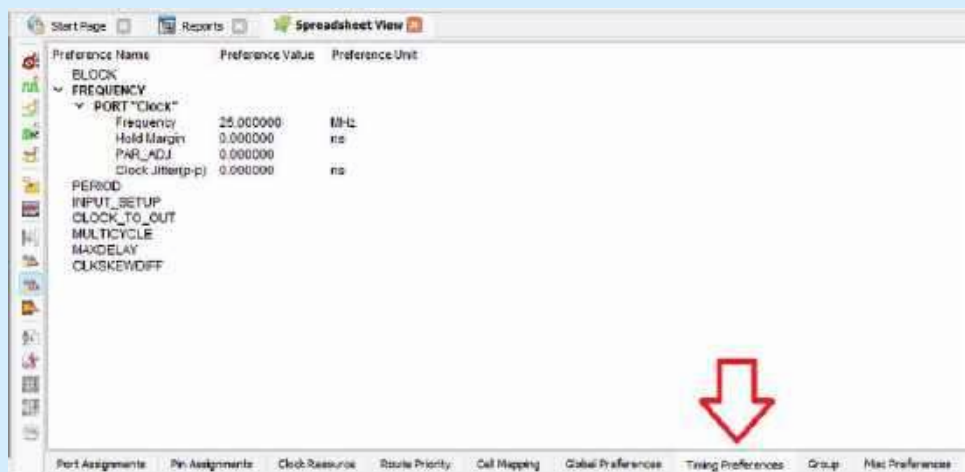
Po wgraniu bitstreamu do FPGA i podłączeniu monitora na ekranie powinien pojawić się obraz podobny do tego z fotografii 1. Grafika powinna być ostra i nieruchoma. Kiedy wejdziemy w ustawiania monitora zobaczymy, że obraz ma rozdzielczość 640×480 px, a częstotliwość odświeżania to 60 Hz, czyli wszystko działa zgodnie z planem (**fotografia 2**).

W ostatnich czasach popularność zyskują proste graficzne wyświetlacze OLED ze sterownikami SSD1309, SSD1306 czy SH1106. Są niedrogie i łatwe w obsłudze. Kontrolery te odbierają dane z mikrokontrolera przez interfejs SPI i wyświetlają ładną grafikę. W następnym odcinku kursu zastosujemy opracowane wcześniej interfejsy SPI oraz VGA, aby stworzyć replikę sterownika OLED. Z punktu widzenia procesora jego obsługa nie będzie różniła się



Name	Group By	Pin	BANK	BANK_VCC	TYPE	ID_TYPE	PULLMODE	DRIVE	SLEWRATE	CLAMP
1.1	All Ports	N/A	N/A	N/A	N/A	N/A	N/A	N/A	N/A	N/A
1.1.1	input	N/A	N/A	N/A	N/A	N/A	N/A	N/A	N/A	N/A
1.1.1.1	Clock	20	3	Auto	N/A	LVCMOS33	NONE	8	N/A	ON
1.1.2	Reset	17	3	Auto	N/A	LVCMOS33	UP	8	N/A	ON
1.2	Output	N/A	N/A	N/A	N/A	N/A	N/A	N/A	N/A	N/A
1.2.1	Blue_o	9	3	Auto	N/A	LVCMOS33	NONE	8	SLOW	OFF
1.2.2	Green_o	10	3	Auto	N/A	LVCMOS33	NONE	8	SLOW	OFF
1.2.3	HSync_o	1	3	Auto	N/A	LVCMOS33	NONE	8	SLOW	OFF
1.2.4	Red_o	78	3	Auto	N/A	LVCMOS33	NONE	8	SLOW	OFF
1.2.5	VSync_o	8	3	Auto	N/A	LVCMOS33	NONE	8	SLOW	OFF

Rysunek 7. Konfiguracja pinów w Spreadsheet



Preference Name	Preference Value	Preference Unit
BLOCK		
FREQUENCY		
PORT "Clock"		
Frequency	25.000000	MHz
Hold Margin	0.000000	ns
PAR_ADJ	0.000000	ns
Clock Jitter(p-p)	0.000000	ns
PERIOD		
INPUT_SETUP		
CLOCK_TO_OUT		
MULTICYCLE		
MAXDELAY		
CLKSKEW		

Rysunek 8. Konfiguracja zegara w Spreadsheet



Fotografia 2. Diagnostyka monitora

niczym od obsługi zwykłego wyświetlacza OLED, ale nasz sterownik będzie pokazywał obraz nie na małej matrycy OLED, lecz na „pełnowymiarowym” monitorze VGA.

Dominik Bieczyński
leonow32@gmail.com

REKLAMA

Mnóstwo doskonałych artykułów, tylko na:

EPcom.pl

koktajl niusów



Oprawa oświetleniowa Linea S Track LED firmy Lena Lighting

Oprawa Linea S Track LED produkcji Lena Lighting powstała z myślą o przestrzeniach użyteczności publicznej. Urządzenie łączy elegancki design z wysoką funkcjonalnością, oferując osiem konfiguracji rozsyłu światła, które dostosowują się do warunków pracy. Producent gwarantuje niską wartość wskaźnika oświecenia (UGR), dzięki czemu oprawa ta nadaje się do pomieszczeń biurowych i przestrzeni coworkingowych. Za sprawą błyskawicznej opcji montażu na szynoprzewodach oprawa oświetleniowa Linea S Track LED stanowi doskonały wybór zarówno do showroomów, lobby czy obiektów handlowych, jak i muzeów, sal wystawowych, a nawet apartamentów. Użytkownik może sam dokonywać zmian aranżacyjnych, nawet po zamontowaniu oprawy Linea S Track LED. Urządzenie stanowi ponadto energooszczędne uzupełnienie projektorów instalowanych na szynoprzewodach. Dostępne są dwie wersje oprawy Linea S Track LED o długości 59 cm oraz 117,5 cm – ze różnicowanymi układami optycznymi, umożliwiającymi dobór właściwego rozwiązania do rozmaitych aplikacji.

<https://tiny.pl/58-q7hqm>



Spiralne przewody sterownicze SIMECH w ofercie TME

Spiralne przewody sterownicze w izolacji poliuretanowej firmy SIMECH przeznaczone są przede wszystkim do urządzeń telekomunikacyjnych oraz robotów i innych maszyn przemysłowych. Nadają się doskonale m.in. do zasilania elementów zamieszczonych na ruchomych głowicach. Przewody są dostępne w dwóch wersjach kolorystycznych: białej oraz czarnej, charakteryzują się dobrą rozciągłością – współczynnik rozciągania jest równy 4:1, co oznacza, że np. spirala o długości 1,5 m może osiągnąć maksymalną długość

6 m, po czym bezpiecznie wraca do swego pierwotnego kształtu. Zewnętrzny płaszcz przewodów wykazuje wysoką odporność m.in. na oleje, mikroby lub hydrolizę, a zakres temperatury pracy (w którym nie dochodzi do trwałej deformacji) rozciąga się od -5 do $+90^{\circ}\text{C}$. Spiralne przewody sterownicze marki SIMECH są dostępne w wersjach o liczbie żył miedzianych od 2 do 7 i przekroju $0,15...1,5\text{ mm}^2$. Opisywane produkty mogą pracować przy napięciu nominalnym 300 lub 500 V i nie rozprzestrzeniają płomienia, co umożliwia ich stosowanie w roli okablowania zasilającego (niektóre wersje przewodów są wyposażone w żyłę PE, wyróżnioną żółto-zieloną barwą izolacji).

<https://tiny.pl/2xfyfg91>



Powiew świeżości w tunelu aerodynamicznym NASA dzięki firmie ABB

Na terenie kampusu NASA w Hampton działa National Transonic Facility (NTF) – jeden z największych na świecie wysokociśnieniowych, kriogenicznych tuneli aerodynamicznych. Dzięki tunelowi możliwe staje się przeprowadzanie złożonych symulacji lotów z prędkością bliską prędkości dźwięku, w temperaturze: od -150 do 65°C . O wyśrubowane parametry obiektu dba szereg urządzeń, w tym wentylator oraz napęd o zmiennej prędkości obrotowej (VSD) – oba dostarczone przez firmę ABB. Po przeszło 20 latach bezproblemowego użytkowania NASA podjęła decyzję o modernizacji dotychczas stosowanych rozwiązań, zamiast instalacji nowych. W chwili uruchomienia (w 1997 roku) współpracujący z tunelem NTF napęd był jednym z lepszych na świecie – działając z mocą 101 MW, pozwalał przesymulować warunki panujące podczas różnego rodzaju lotów. Przed przystąpieniem do modernizacji tunelu specjaliści ABB ocenili poziom wydajności i stan połączeń mechanicznych napędu. Korzystając z najnowocześniejszych podzespołów energoelektronicznych, wymieniono m.in. jednostkę sterującą napędu (a ściślej rzecz ujmując – niewielką część falownika), ograniczając do minimum ilość odpadów, a także czas potrzebny na prace serwisowe. Starano się też dopasować całkowicie nowe układy do istniejącej instalacji napędu NTF, dzięki czemu całość funkcjonuje z taką samą mocą, jak przed modernizacją, lecz przy jeszcze wyższym poziomie niezawodności. Dzięki przeprowadzonym pracom przedłużono, o minimum 10 lat, żywotność tego napędu – i być może pierwszy od dłuższego czasu prace nad nowym samolotem, na dodatek ponaddzwiękowym, będą miały swój początek właśnie w NTF.

<https://tiny.pl/7wrkmm27>



Hybrydowe kondensatory elektrolityczne z serii ZL firmy Panasonic do aplikacji motoryzacyjnych

Panasonic zaprezentował nowe kondensatory elektrolityczne z serii ZL, których konstrukcja bazuje na autorskim elektrolicie pozwalającym uzyskiwać wysoką pojemność przy zachowaniu niewielkich wymiarów. Opisywane elementy charakteryzuje przede wszystkim długa żywotność (przekraczająca 4 tysiące godzin), którą dało się osiągnąć dzięki specjalnej technologii formowania polimerów, wchodzących w skład elektrolitu. Kondensatory elektrolityczne z serii ZL dostępne są w 5 rozmiarach i pod względem pojemności przewyższają serię kondensatorów ZC, należącą również do portfolio marki Panasonic. Szeroki zakres temperatury pracy (od -55 do $+135^{\circ}\text{C}$) jest wystarczający do większości zastosowań. Wartość dopuszczalnego prądu tętnień waha się w przedziale $0,6...4,4$ A (przy częstotliwości do 100 kHz i temperaturze $125^{\circ}\text{C}...135^{\circ}\text{C}$). Elementy są przeznaczone m.in. do zastosowań w przetwornicach DC/DC, falownikach, systemach ADAS, jednostkach kontroli strefy (ZCU) oraz systemach ECU.

<https://news.panasonic.com/global/press/en240228-2>



Innowacyjna technologia samochodowa LG Electronics

Jazda samochodem może być niełatwa, zwłaszcza na skomplikowanych drogach bądź nieznanach trasach, gdzie oszacowanie odległości, a także właściwe rozpoznanie przez kierowcę wskazówek nawigacji staje się dość kłopotliwe. Najnowsza technologia samochodowa AR firmy LG Electronics dostarcza wyraźne linie prowadzące, za którymi można z łatwością podążać. Oferowana przez LG Electronics technologia korzysta z zaawansowanych algorytmów, które umożliwiają integrację danych m.in. z czujników auta, jego kamer i systemów nawigacji. Dzięki informacjom płynącym z systemów wspomagania kierowcy (ADAS), z użyciem komunikacji w standardzie V2X, technologia ta pozwala przewidywać przeróżne, potencjalnie niebezpieczne dla kierowców zagrożenia, w oparciu o lokalizację pojazdu i otaczające środowisko. Zasadniczą cechą opisywanej technologii jest pokaz dynamicznych obrazów AR, dzięki którym prowadzenie auta staje się bardziej intuicyjne

– zwłaszcza dzięki temu, że symbole pojazdu, a także kierunku jazdy (np. w formie strzałek) dopasowują się do warunków na drodze niemal błyskawicznie. Grafika AR może się płynnie rozdzielać, zmieniać kształt i łączyć, przekazując informacje o m.in. prędkości auta, nachyleniu terenu, możliwej zmianie pasa ruchu czy nawet samej trasie – w przejrzysty i zrozumiały sposób. Firma LG Electronics planuje wykorzystać tę technologię, by opracować nowe rozwiązania dla czołowych producentów aut, przyspieszając w ten sposób rozwój innowacyjnego oprogramowania i generując przychody z tytułu tantiem, które są związane z opatentowaną technologią.

https://tiny.pl/b_d5ty0j



Najnowsza seria elektronicznych wyświetlaczy papierowych EPD firmy WINSTAR

W odróżnieniu od klasycznych rozwiązań (np. LCD czy OLED), ekrany EPD podtrzymują wyświetlony obraz przez dowolnie długi czas, przy czym pobór mocy ma miejsce tylko podczas aktualizacji prezentowanej treści. Najnowsze wyświetlacze firmy WINSTAR oferują – oprócz energooszczędności – także szerokie kąty widzenia do 180° i doskonałą czytelność w blasku słońca. Niektóre modele wyposażone są w panele dotykowe i/lub oświetlenie adaptacyjne (umieszczone z przodu), co zwiększa funkcjonalność prezentowanych rozwiązań, zapewniając wyjątkowo spójne wrażenia wizualne. W portfolio marki WINSTAR można znaleźć ekrany EPD czarno-białe, trójkolorowe oraz czterokolorowe, a dodatkowo producent oferuje możliwość adaptacji wyświetlaczy do konkretnych potrzeb, np. na drodze konfiguracji FPC czy integracji z dodatkowymi elementami mechanicznymi.

<https://www.winstar.com.tw/pl/news/detail/237.html>



Wyjątkowe kolumny aktywne 4305P firmy JBL

Aktywne kolumny JBL 4305P to niesłychanie ciekawa propozycja dla każdego miłośnika muzyki, pozwalająca na niezawodne odtwarzanie wysokiej jakości dźwięku przesyłanego strumieniowo. Istotną w tym zasługą zarówno wbudowanego, jednocalowego głośnika kompresyjnego, jak przetwornika cyfrowo-analogowego o rozdzielczości 24 bitów i częstotliwości próbkowania 192 kHz. Kolumny 4305P, o mocy wyjściowej 300 W, wyposażone są w takie rodzaje wejść jak: XLR/6,3 mm, TRS, jack 3,5 mm,

Toslink oraz USB. Bezprzewodową transmisję dźwięku umożliwia łączy Bluetooth, a w standardzie jest także obsługa takich rozwiązań jak Google Chromecast czy Apple AirPlay 2. Urządzenia mogą być sterowane za pomocą znajdującego się w zestawie pilota lub aplikacji MusicLife. Kolumny aktywne 4305P produkcji JBL produkowane są w obudowach z wykończeniem fornirowym klasy meblowej, w kolorze orzecha naturalnego lub orzecha czarnego.

<https://tiny.pl/pfq2v3p1>

Nowe przyrządy sieciowe do monitorowania natężenia pola elektromagnetycznego

Smog elektromagnetyczny, wynikający z lawinowego wzrostu liczby urządzeń IoT komunikujących się w popularnych pasmach ISM, stanowi istotny problem zarówno z punktu widzenia bezpieczeństwa, jak i użyteczności infrastruktury bezprzewodowej. MonitEM-IoT to nowoczesne rozwiązanie opracowane przez markę Wavecontrol, którego celem jest monitorowanie pól elektromagnetycznych w nowoczesnych miastach i infrastrukturze przemysłu 4.0. W ofercie producenta znalazły się dwie wersje urządzenia:



- MonitEM-IoT-N-RF8 (pasmo 500 kHz...8 GHz, zasilanie bateryjne o czasie pracy powyżej 3 lat, komunikacja bezprzewodowa LTE-M/Nb-IoT),
- MonitEM-IoT-N-RF8 (pasmo 5 MHz...8 GHz, zasilanie PoE, komunikacja przez łączy Ethernet).

Pozostałe parametry są wspólne dla obydwu wariantów:

- Technologia pomiaru: 3-osiowy, izotropowy czujnik pola elektrycznego RMS na bazie diod,
- Charakterystyka częstotliwościowa: płaska ($\pm 1,5$ dB w paśmie od 1 MHz do 6 GHz),
- Zakres pomiarowy: 0,7...250 V/m,
- Częstość pomiarów: konfigurowalna,
- Uśrednianie wyników: 6 min/30 min (wartość konfigurowalna),
- Przechowywanie danych pomiarowych: >1 miesiąca w przypadku utraty komunikacji,
- Programowalne alarmy: poziom natężenia pola, wyczerpanie baterii, błąd komunikacji,
- Obsługa NFC podczas konfiguracji: tak,
- Zawartość zestawu: mocowanie do ściany, masztu lub statywu,
- Zgodność z normami: ITU-T K83, ICNIRP2020,
- Zakres temperatur pracy: -20...+50°C,
- Wymiary: 24,0×7,5×6,0 cm,
- Masa: 550 g,
- Stopień ochrony: IP65.

<https://www.wavecontrol.com>

Zoptymalizowana wersja bezsmarowych łożysk igubal ESQM firmy igus do śledzenia modułów słonecznych

Od blisko 15 lat łożyska igubal ESQM stosowane są w konstrukcji ruchomych modułów fotowoltaicznych, automatycznie ustawiających się w kierunku słońca. Teraz nadeszła pora na oficjalną demonstrację ich zoptymalizowanej wersji – igubal ESQM 2.0 – do budowy której zastosowano najnowszą generację tworzywa sztuczne. Do łożyska przymocowany jest kwadratowy profil, który pełni funkcję wspornika modułów PV i współpracuje z dzieloną oprawą oraz dwoma sferycznymi półpowłokami, które ułatwiają montaż. Od teraz półpowłoki mają kształt litery U, a zintegrowana szczelina umożliwia wyprowadzenie kabli także przez światło łożyska. Za sprawą



dwóch niewielkich bolców, które są umieszczone pod oprawą, zapewniona jest prosta i bezślizgowa instalacja, a smukła konstrukcja przekłada się na dodatkową przestrzeń montażową. W ofercie producenta jest też dostępna oryginalna poprzeczka rozszerzająca z metalu, pozwalająca bezpiecznie podtrzymywać czaszą kulistą w oprawie. Zastosowanie bezsmarowych komponentów redukuje koszty konserwacji, a także redukuje ryzyko przedostania się środka smarnego do gleby na otwartych przestrzeniach. Wszystkie opisane rozwiązania przekładają się na niezwykle płynną pracę łożyska igubal ESQM 2.0 – produktu, który z łatwością pracuje na sucho, co stanowi jego najważniejszy atut.

<https://press.igus.pl/bezsmarowe-łożyska-na-słonce/>



Modernizacja stacji bazowych (BTS) na przykładzie firmy Orange

Prace modernizacyjne – obok uruchamiania nowych stacji bazowych – przyczyniają się do wzrostu pojemności czy też wzmocnienia zasięgu sieci komórkowej. Tylko od początku 2024 roku firma Orange zmodernizowała 1670 swych stacji, dokładając na wielu z nich aż 820 warstw sieci LTE. Podczas modernizacji klasycznej stacji bazowej (BTS) wymienia się przede wszystkim anteny razem z modułami radiowymi i sterującymi, po czym ma miejsce aktualizacja oprogramowania. I chociaż w wielu przypadkach oznacza to chwilowy brak dostępu do usług sieci komórkowej oraz wiąże się ze zmianami dostawców sprzętu telekomunikacyjnego, to w ostatecznym rozrachunku warto to jednak przeboleć – zwłaszcza że po modernizacji zapewniane są dodatkowe warstwy sieci LTE w 4 pasmach częstotliwości: 800, 1800, 2100 i 2600 MHz. O tym, które dokładnie pasma należy uruchomić, decydują przede wszystkim lokalne potrzeby sieciowe, lokalizacja stacji i topografia terenu. Modernizując stacje bazowe, dąży się do przygotowania do wyłączenia sieci 3G oraz refarmingu, czyli uruchomienia LTE w paśmie 900 MHz, a przyszłościowo także 5G w paśmie C. Nowy sprzęt daje możliwość wykorzystania nowszych technologii i generacji telefonii komórkowej, co pozwala znacząco zwiększyć zasięg sieci.

Jakub Tyburski
jakub.tyburski@elportal.pl

Zamek szyfrowy do kontroli dostępu

Zaprezentowane urządzenie można zastosować jako element prostego systemu kontroli dostępu, w którym tylko osoba znająca kod cyfrowy może włączyć jakieś urządzenie, odblokować system alarmowy, otworzyć drzwi itp. Układ można skonfigurować jako włącznik bistabilny lub monostabilny z programowalnym czasem stanu aktywnego. Pomimo zastosowania klawiatury numerycznej o 12 przyciskach i architekturze matrycowej 3×4, układ udało się zbudować w oparciu na... 8-pinowym mikrokontrolerze PIC12F675!

Symetryzator – odwracacz fazy audio

Moduł został opracowany w celu uzupełnienia układów współczesnych, cyfrowych końcówek mocy, np.: TPA3251, TPA3116, wymagających – dla osiągnięcia możliwie najwyższych parametrów użytkowych – sterowania ze źródła sygnału symetrycznego, najlepiej o niewielkiej rezystancji wyjściowej. Urządzenie można oczywiście zastosować także do połączenia komercyjnego sprzętu audio (wyposażonego w wyjścia niesymetryczne) z profesjonalnymi kartami muzycznymi, mikserami lub wzmacniaczami, w których dostępne są tylko wejścia zbalansowane.



Mikrokontrolery i procesory aplikacyjne

Mikrokontrolery to niezwykle bogata i różnorodna, a jednocześnie bardzo ważna gama elementów elektronicznych. Właściwy dobór układu do projektowanej aplikacji wpływa nie tylko na funkcjonalność, energooszczędność, niezawodność końcowego urządzenia, ale też na bezpieczeństwo danych, które urządzenie gromadzi, przechowuje, przetwarza i transmituje. W artykule przypomnimy najważniejsze pojęcia dotyczące nowoczesnych mikrokontrolerów oraz opiszemy wybrane kategorie tych układów.



Oscyloskopy cyfrowe

Oscyloskopy cyfrowe dawno wyparły już z rynku swoich analogowych krewnych i niepodzielnie rządzą we wszystkich pracowniach elektronicznych – serwisach, prototypowniach działów R&D, laboratoriach uczelnianych czy nawet montażowniach. Trzecia dekada XXI wieku przyniosła ogromne zmiany na rynku tych fundamentalnych urządzeń pomiarowych,

standardem stały się bowiem rozwiązania 12-bitowe: dawniej tego typu oscyloskopy należały do ścisłej czołówki najdroższej aparatury wysokiej klasy, dziś dwanaście bitów to już codzienność także w segmencie entry-level. Jakże jeszcze trendy są zauważalne na rynku DSO i MSO i czego można się spodziewać po nadchodzących latach? W lutowym artykule z serii „Elektronika w Praktyce” uważnie (i krytycznie!) przyglądamy się właśnie tej grupie urządzeń.



Wykaz firm ogłaszających się w tym numerze „Elektroniki Praktycznej”

AKSOTRONIK	27
ARTRONIC	62
BORNICO	69
COMPUTER CONTROLS	7
ELTY.PL	65
FERYSTER	9
HAMMOND	5
LASTENIC LASER	21
NEWLAND	81
SEMICON	19

Miesięcznik „Elektronika Praktyczna” (12 numerów w roku) jest wydawany przez AVT Korporacja Sp. z o.o. we współpracy z wieloma redakcjami zagranicznymi.



Wydawnictwo:
AVT Korporacja Sp. z o.o.
03-197 Warszawa, ul. Leszczyńska 11
tel. 22 257 84 99, e-mail: avt@avt.pl

Wydawca:
Wiesław Marciniak

Adres redakcji:
03-197 Warszawa, ul. Leszczyńska 11
e-mail: redakcja@ep.com.pl, www.ep.com.pl

Redaktor Naczelny:
Przemysław Musz

Redaktor Programowy, Przewodniczący Rady Programowej:
Piotr Zbysiński

Menedżer Magazynu:
Katarzyna Gugala, tel. 22 257 84 64

Szef Pracowni Konstrukcyjnej:
Jakub Sobański

Zespół marketingu i reklamy:
Katarzyna Gugala, Bożena Krzykawska,
Grzegorz Krzykowski, Grzegorz Lalak

Stali współpracownicy:
Lucjan Bryndza, Nikodem Czechowski, Jarosław Doliński,
Andrzej Gawryluk, Krzysztof Górski, Tomasz Jabłoński,
Paweł Kowalczyk, Henryk Kowalski, Rafał Kozik,
Michał Kurzela, Jakub Nowicki, Szymon Panecki,
Adam Sobczyk, Damian Sosnowski, Ryszard Szymaniak,
Adam Tatuś, Jakub Tyburski, Robert Wołgajew

Uwaga!
Kontakt z wymienionymi osobami jest możliwy via e-mail, według schematu: imię.nazwisko@ep.com.pl

DTP, okładka, redakcja strony internetowej www.ep.com.pl:
MAD Sp. z o.o.

Prenumerata w Wydawnictwie AVT
www.ulubionykiosk.pl lub tel. 22 257 84 22
(godz. 10.00–14.00)
e-mail: prenumerata@avt.pl

Prenumerata w RUCH S.A.
tel. 801 800 803, 22 717 59 59
e-mail: prenumerata@ruch.com.pl

Copyright AVTKorporacja Sp. z o.o.
03-197 Warszawa, ul. Leszczyńska 11

Projekty publikowane w „Elektronice Praktycznej”. mogą być wykorzystywane wyłącznie do własnych potrzeb. Korzystanie z tych projektów do innych celów, zwłaszcza do działalności zarobkowej, wymaga zgody redakcji „Elektroniki Praktycznej”. Przedruk oraz umieszczanie na stronach internetowych całości lub fragmentów publikacji zamieszczanych w „Elektronice Praktycznej”. jest dozwolone wyłącznie po uzyskaniu zgody redakcji. Redakcja nie odpowiada za treść reklam i ogłoszeń zamieszczanych w „Elektronice Praktycznej”.



Ulubiony Kiosk

Pobierz bezpłatnie multimedialne dodatki do tego wydania Elektroniki Praktycznej

**Projekty, miniprojekty, materiały do
artykułów i kursów oraz wiele innych!**

*** Kupiłeś magazyn
w Ulubionym
Kiosku lub masz
prenumeratę?
Multimedialne dodatki
będą odblokowane
automatycznie!**

*** Zakupiłeś czasopismo
u zewnętrznego
dystrybutora?
Odblokuj bibliotekę
multimediów
samodzielnie.**

Szczegóły na UlubionyKiosk.pl/media